
Einführung in Web- und Data-Science

Prof. Dr. Ralf Möller

Universität zu Lübeck

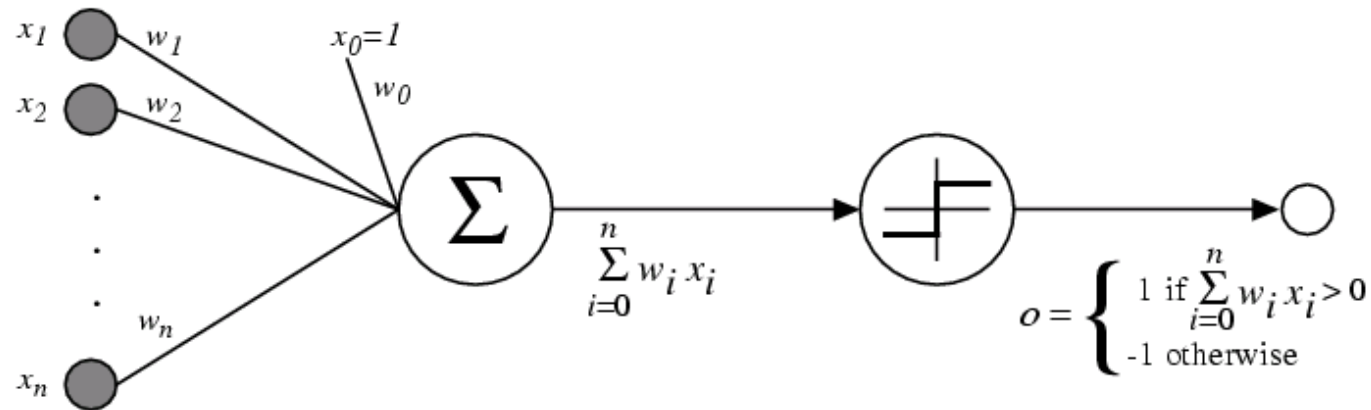
Institut für Informationssysteme

Repräsentation von Funktionen ...

- ... durch Perzeptrons
 - Ein-Ebenen-Netzwerk (Linearer Klassifikator)
 - Mehrebenen-Netzwerke
 - Lernregel Fehlerrückführung (Backpropagation)

Frank Rosenblatt, The Perceptron--a perceiving and recognizing automaton. Report 85-460-1, Cornell Aeronautical Laboratory, **1957**

Perzeptron

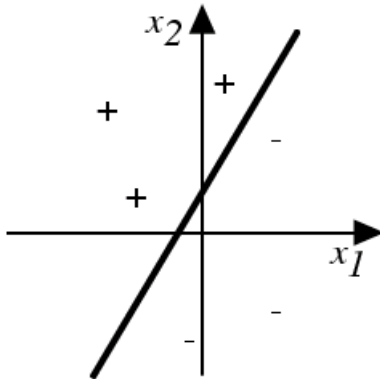


$$o(x_1, \dots, x_n) = \begin{cases} 1 & \text{if } w_0 + w_1 x_1 + \dots + w_n x_n > 0 \\ -1 & \text{sonst} \end{cases}$$

Manchmal einfacher geschrieben als (Ann.: $x_0 = 1$):

$$o(\vec{x}) = \begin{cases} 1 & \text{if } \vec{w} \cdot \vec{x} > 0 \\ -1 & \text{sonst} \end{cases}$$

Entscheidungslinie eines Perzeptrons



Repräsentiert lineare Funktion $y = mx + b$

- Was machen die Gewichte?

$$g(x_1, x_2) = AND(x_1, x_2)?$$

Verallgemeinerung auf Entscheidungsebenen möglich

Vgl.: Warren McCulloch und William Pitts: *A logical calculus of the ideas immanent in nervous activity*. In: *Bulletin of Mathematical Biophysics*, Bd. 5, S. 115–133, **1943**

Trennlinien

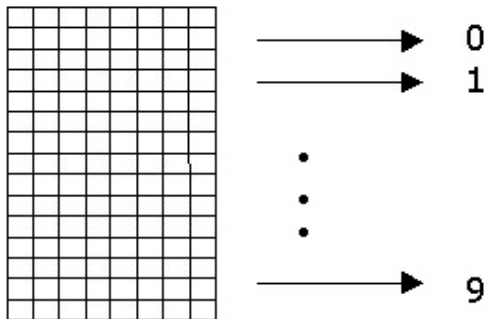
- Wenn $w_2x_2 + w_1x_1 + w_0 > 0$, dann wird für (x_1, x_2) ein (+) vergeben.
- $w_2x_2 + w_1x_1 + w_0 = 0$ ist Trennlinie. Warum?
- Es wird die Gerade $x_2 = -w_1/w_2 x_1 - w_0/w_2$ als Trennlinie verwendet.
- Wir betrachten $y = mx + b$
- Bedingung lautet: $y - mx - b > 0$
- Wähle $m=1$ und b positiv. Dann betrachte $(x, y) = (0, 2b)$. Das liegt über der Geraden.
- $2b - 0 - b > 0$
- Also wird $(0, 2b)$ mit 1 bewertet (+)

Lassen sich alle Funktionen repräsentieren?

- Kann die Fehlerfunktion immer sinnvoll minimiert werden?
- XOR-Problem
 - Einführung weiterer Dimensionen
 - Erweiterung der Daten?
 - Besser: Einführung weiterer Ebenen im Netz

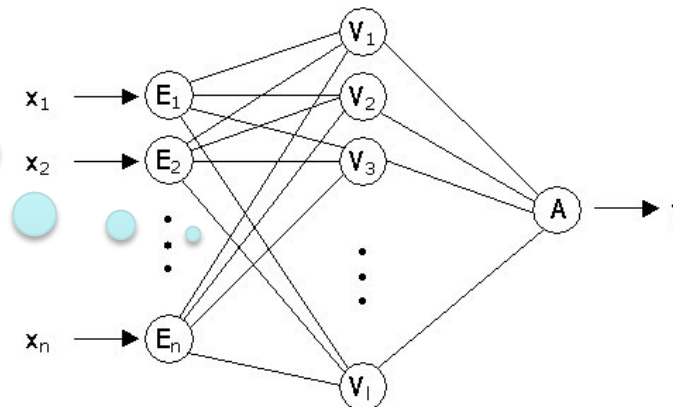
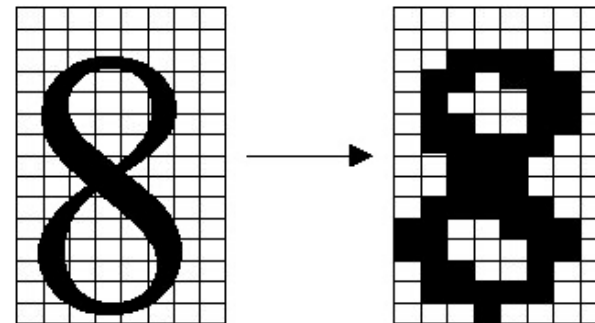
Ein Anwendungsbeispiel

Das folgende Netz soll Ziffern von 0 bis 9 erkennen. Dafür wird zunächst das *Eingabefeld* in 8x15 Elemente aufgeteilt:

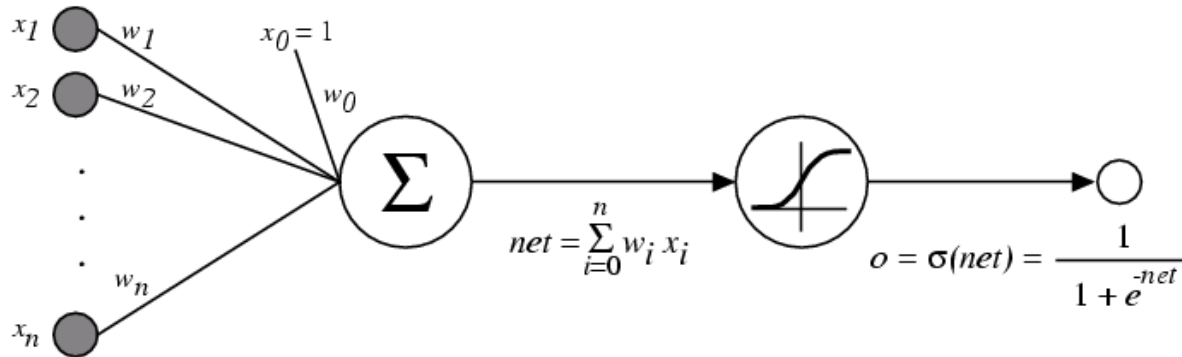


Wie können wir die Parameter automatisch so bestimmen, dass geschriebene Ziffern erkannt und als Ausgabe y ausgegeben werden?

Die geschriebene Ziffer wird in eine Folge von Nullen und Einsen umgewandelt, wobei 0 für leere und 1 für übermalte Rasterpunkte steht:



Verwendung kontinuierlicher Funktionen



$\sigma(x)$ ist die Sigmoid-Funktion

$$\frac{1}{1 + e^{-x}}$$

-0.06

W1

-2.5

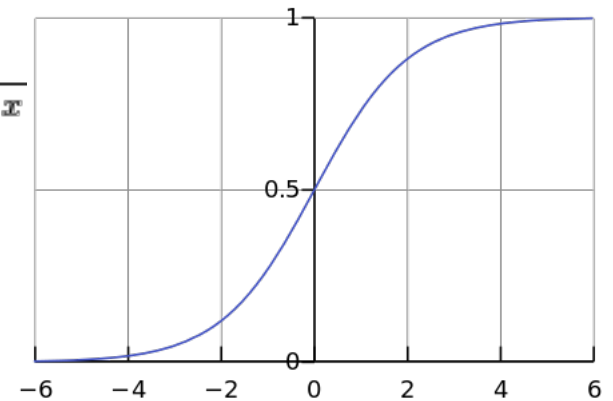
W2

W3

1.4

$f(x)$

$$f(x) = \frac{1}{1 + e^{-x}}$$



-0.06

2.7

-2.5

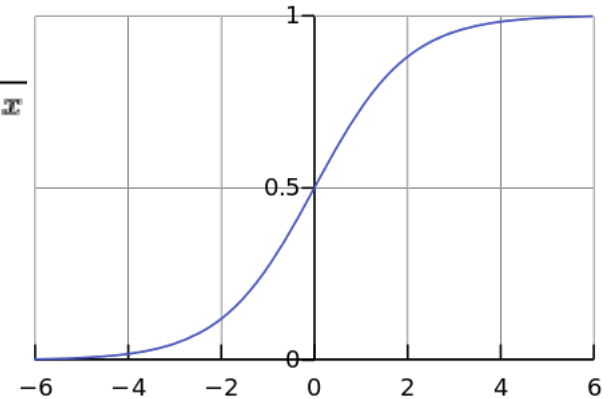
-8.6

0.002

1.4

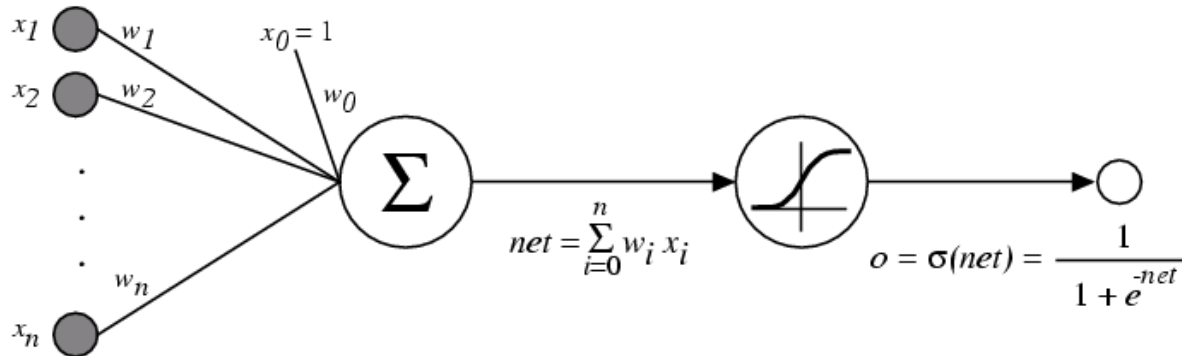
$f(x)$

$$f(x) = \frac{1}{1 + e^{-x}}$$



$$x = -0.06 \times 2.7 + 2.5 \times 8.6 + 1.4 \times 0.002 = 21.34$$

Sigmoid-Einheit



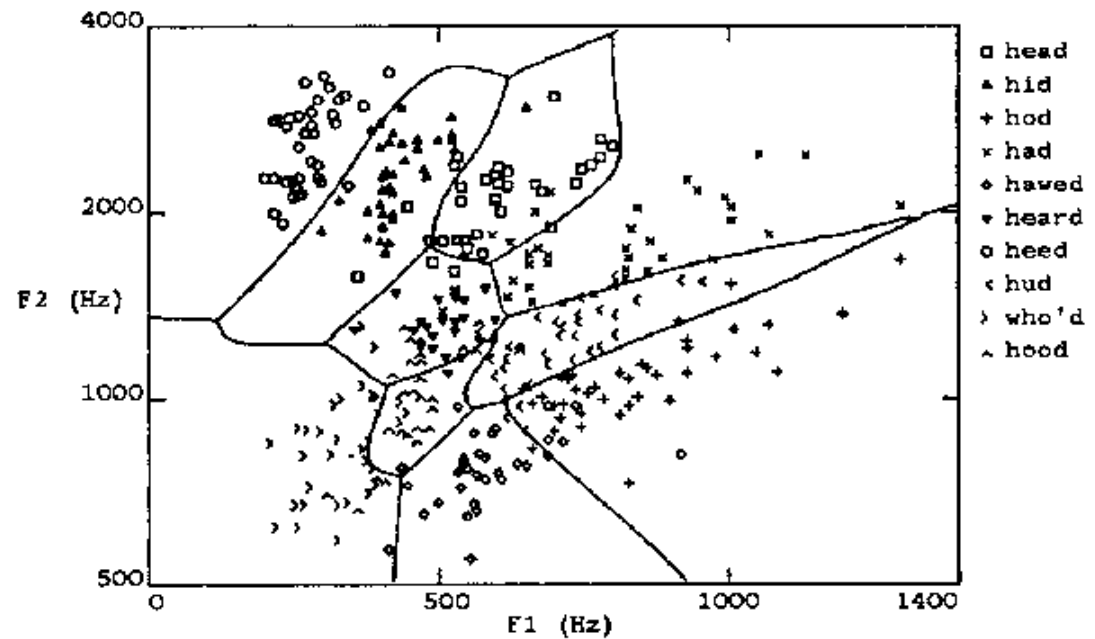
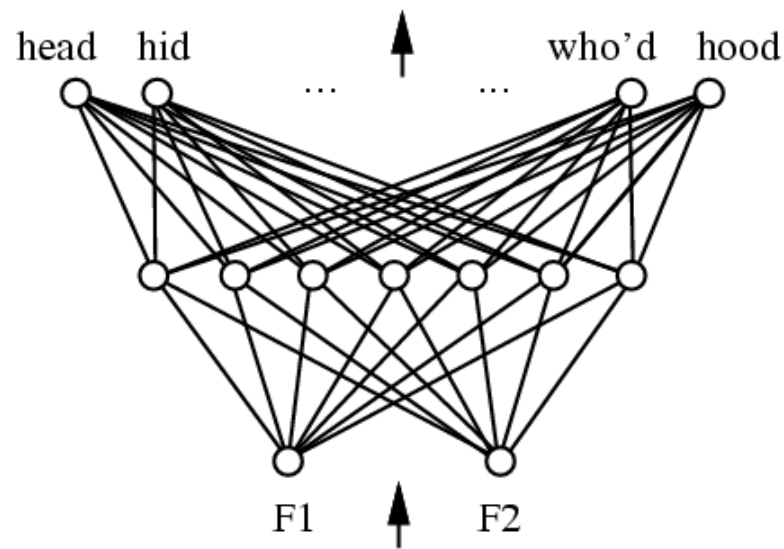
$\sigma(x)$ ist die Sigmoid-Funktion (auch: logistische Funktion)

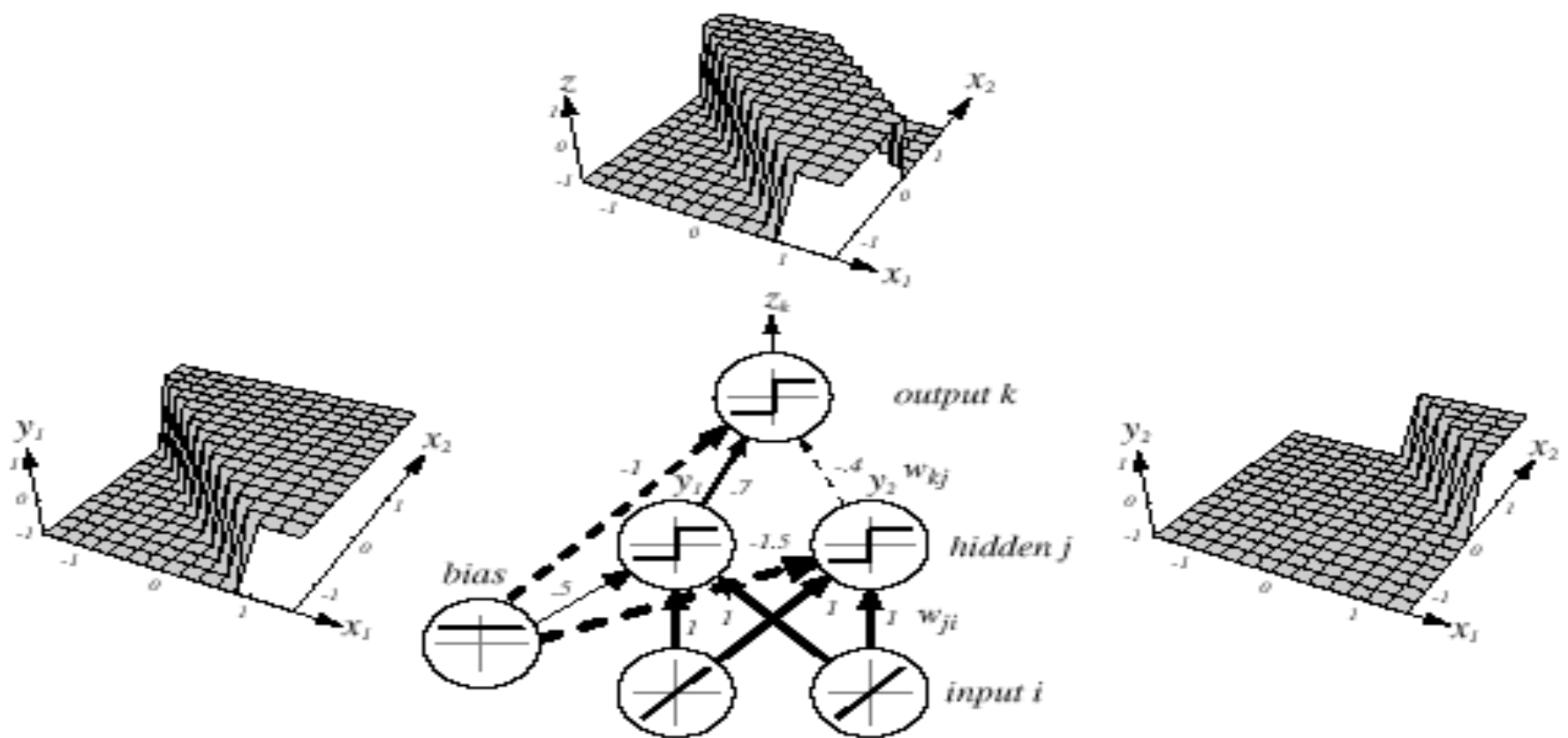
$$\frac{1}{1 + e^{-x}}$$

Eigenschaft: $\frac{d\sigma(x)}{dx} = \sigma(x)(1 - \sigma(x))$ (Gradient)

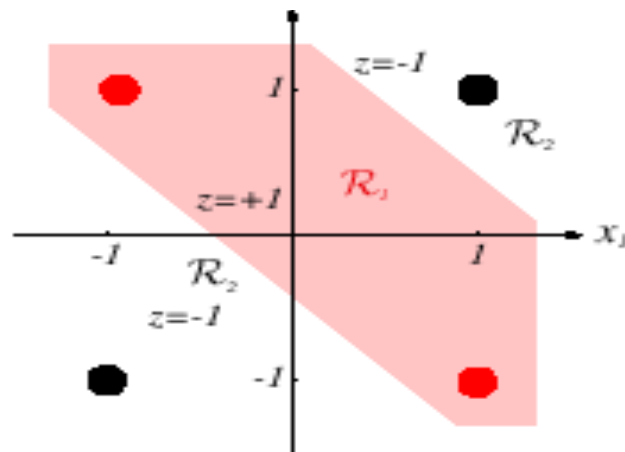
- Wir können den Gradienten verwenden, um die Einheit anzupassen
- Wir kommen gleich darauf zurück

Mehr-Ebenen Netze von Sigmoid-Einheiten





$$Z = x_1 \text{ XOR } x_2 = (x_1 \text{ OR } x_2) \text{ AND NOT}(x_1 \text{ AND } x_2)$$



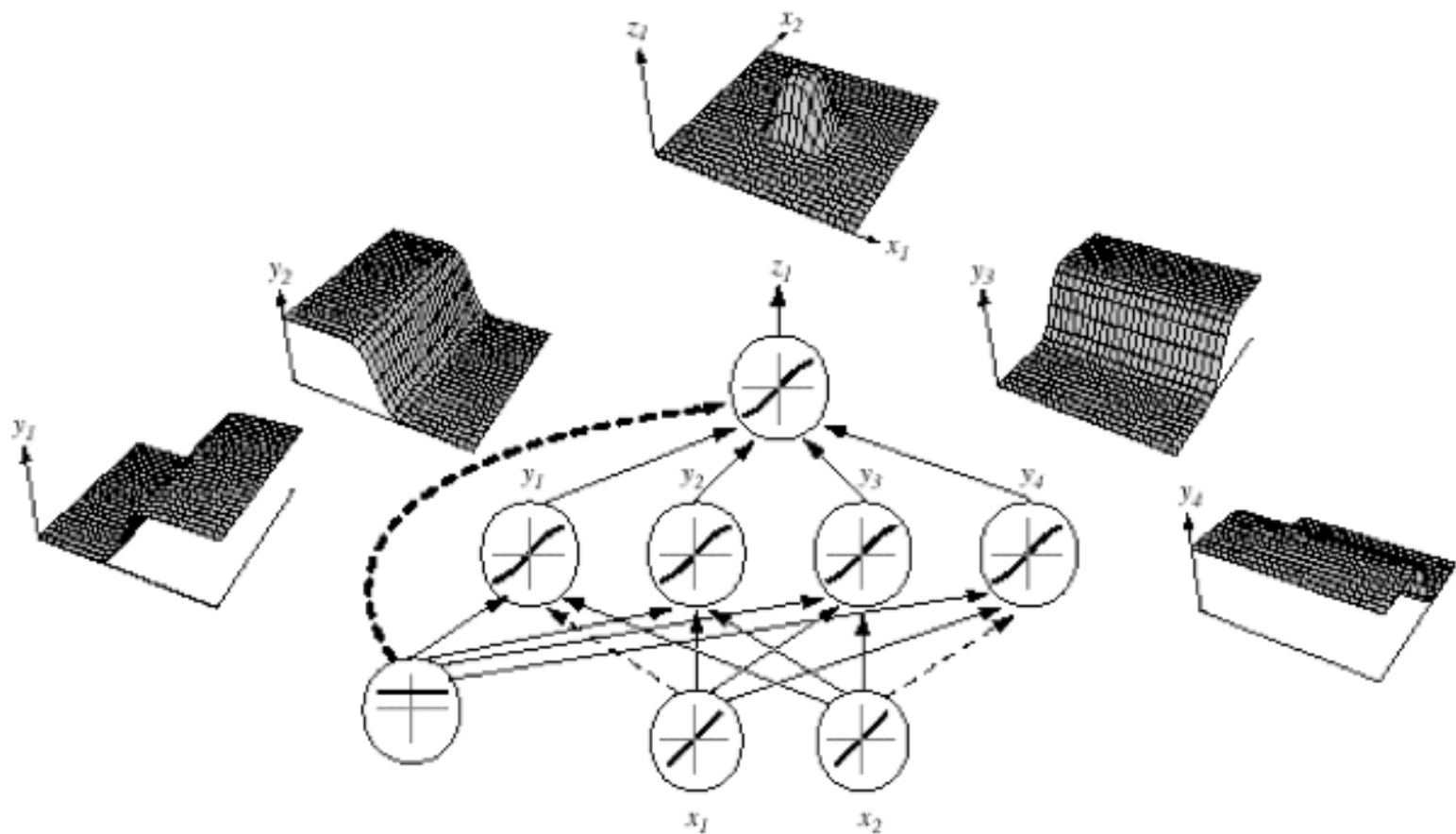


FIGURE 6.2. A 2-4-1 network (with bias) along with the response functions at different units; each hidden output unit has sigmoidal activation function $f(\cdot)$. In the case shown, the hidden unit outputs are paired in opposition thereby producing a “bump” at the output unit. Given a sufficiently large number of hidden units, any continuous function from input to output can be approximated arbitrarily well by such a network. From: Richard O. Duda, Peter E. Hart, and David G. Stork, *Pattern Classification*. Copyright © 2001 by John Wiley & Sons, Inc.

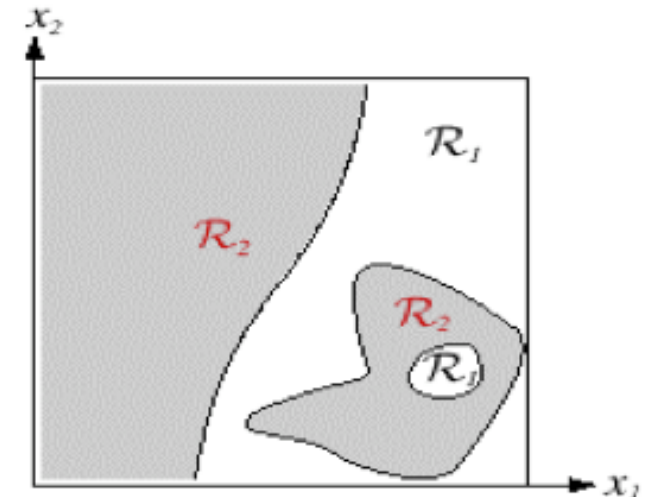
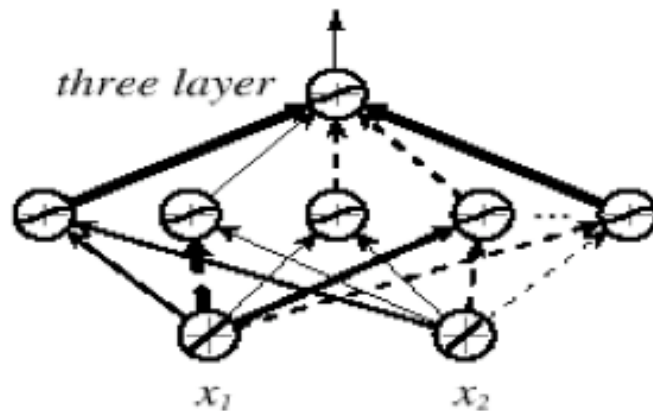
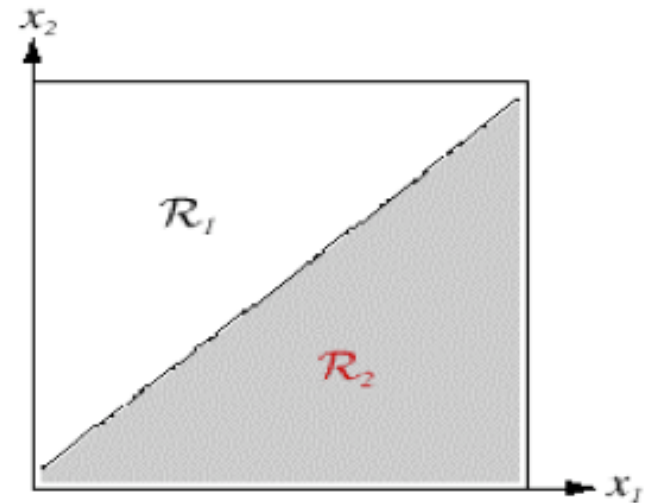
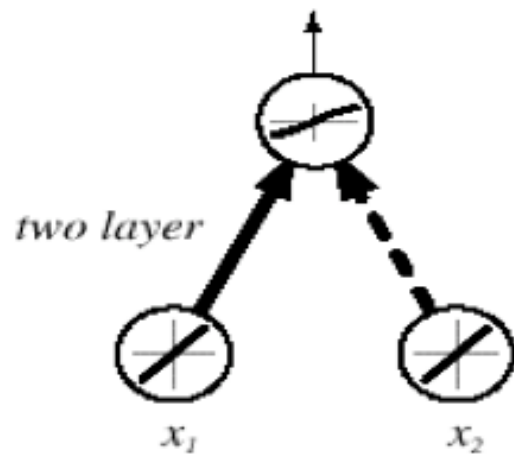
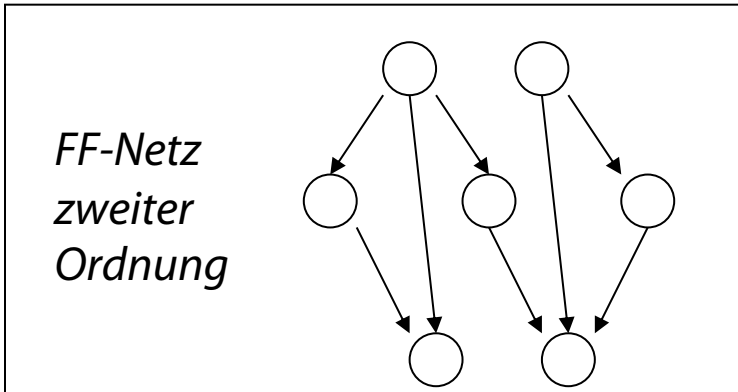
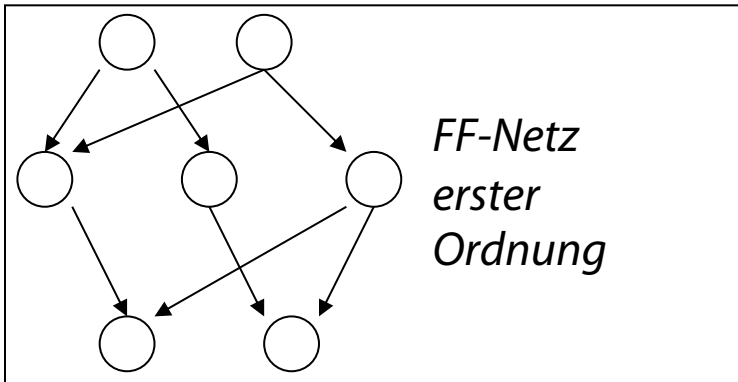


FIGURE 6.3. Whereas a two-layer network classifier can only implement a linear decision boundary, given an adequate number of hidden units, three-, four- and higher-layer networks can implement arbitrary decision boundaries. The decision regions need not be convex or simply connected. From: Richard O. Duda, Peter E. Hart, and David G. Stork, *Pattern Classification*. Copyright © 2001 by John Wiley & Sons, Inc.

Netztopologien

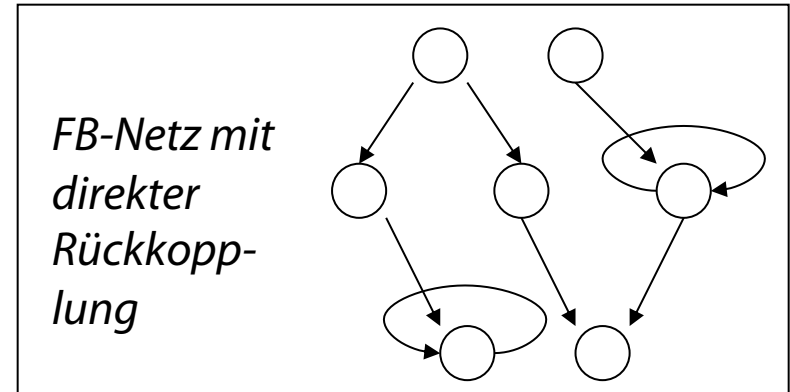
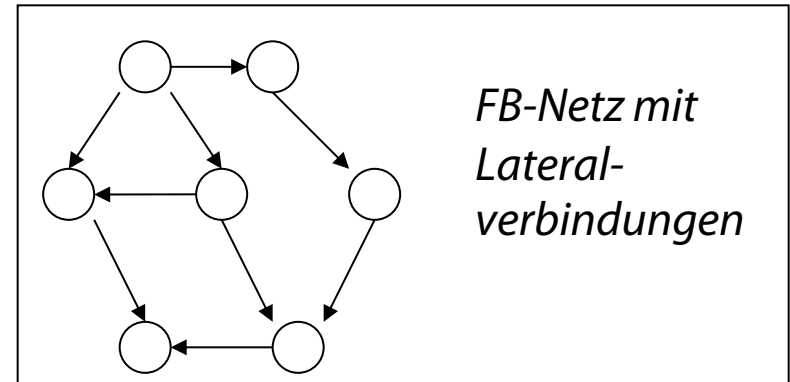
FeedForward-Netze:

Gerichtete Verbindungen nur von niedrigen zu höheren Schichten



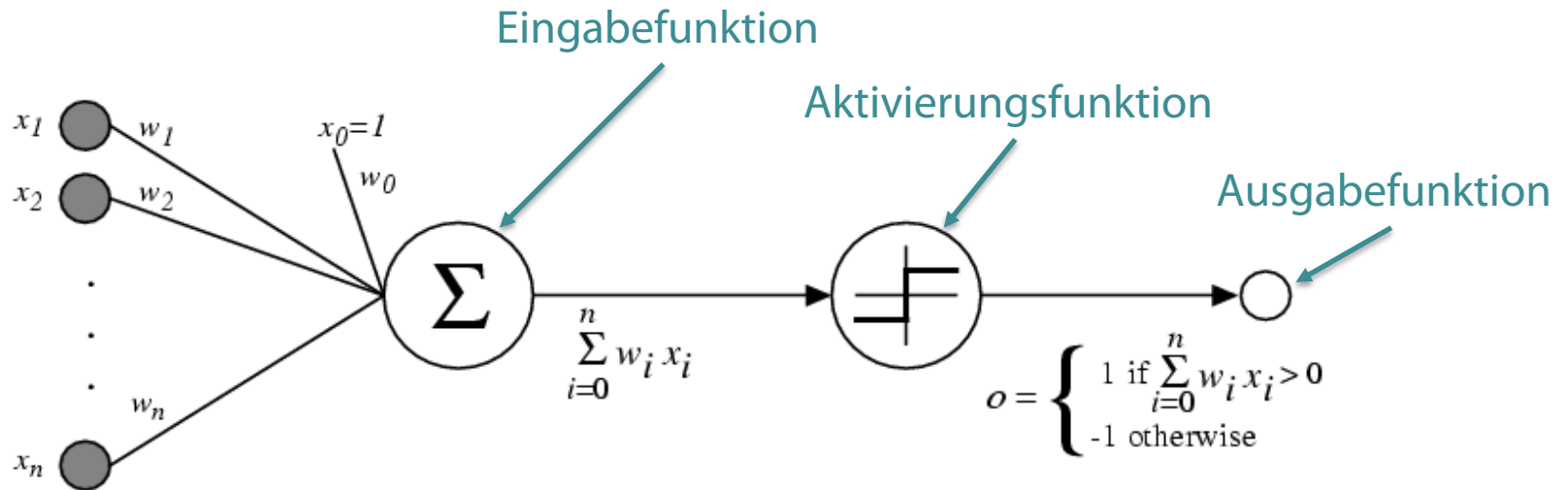
FeedBack-Netze (rekurrente Netze):

Verbindungen zwischen allen Schichten möglich (Abrollen)



Netzwerk von Perzeptrons

- Perzeptron:



- Was muss ich vor dem Lernen/Trainieren bestimmen?
 - Netzwerk-Topologie?
 - Perzeptron bezogen?

Die Eingabefunktion

Die Eingabe- oder Propagierungsfunktion berechnet aus dem Eingabevektor $\vec{x} = (x_1, \dots, x_m)$ und dem Gewichtsvektor $\vec{w}_k = (w_{1,k}, \dots, w_{m,k})$ den Nettoinput des k -ten Knotens.

Es gibt folgende Inputfunktionen:

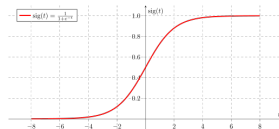
- Summe: $net_k = f_{in}(\vec{w}_k, \vec{x}) = \sum_{i=1}^m w_{i,k} \cdot x_i$
- Maximalwert: $net_k = f_{in}(\vec{w}_k, \vec{x}) = \max_i (w_{i,k} \cdot x_i)$
- Produkt: $net_k = f_{in}(\vec{w}_k, \vec{x}) = \prod_{i=1}^m w_{i,k} \cdot x_i$
- Minimalwert: $net_k = f_{in}(\vec{w}_k, \vec{x}) = \min_i (w_{i,k} \cdot x_i)$

Die Aktivierungsfunktion

Mit der Aktivierungsfunktion (auch: Transferfunktion) wird aus dem Nettoinput net_k der Aktivierungszustand a_k eines Knotens berechnet. Folgende Aktivierungsfunktionen sind gebräuchlich:

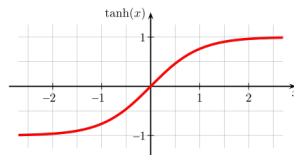
- Lineare Aktivierungsfkt.: $a_k = f_{act}(net_k) = c_k \cdot net_k$
- Binäre Schwellenwertfkt.: $a_k = f_{act}(net_k) = \begin{cases} 1, & \text{falls } net_k \geq \theta_k, \\ 0, & \text{sonst.} \end{cases}$
- Fermi-Fkt. (logistische Fkt.): $a_k = f_{act}(net_k) = \frac{1}{1+e^{-\frac{net_k}{T}}}$

Schwellenwert,
häufig 0



Spezialfall: Sigmoid: $T=1$

- Tangens hyperbolicus: $a_k = f_{act}(net_k) = \frac{e^{net_k} - e^{-net_k}}{e^{net_k} + e^{-net_k}} = \frac{1 + \tanh(net_k)}{2}$



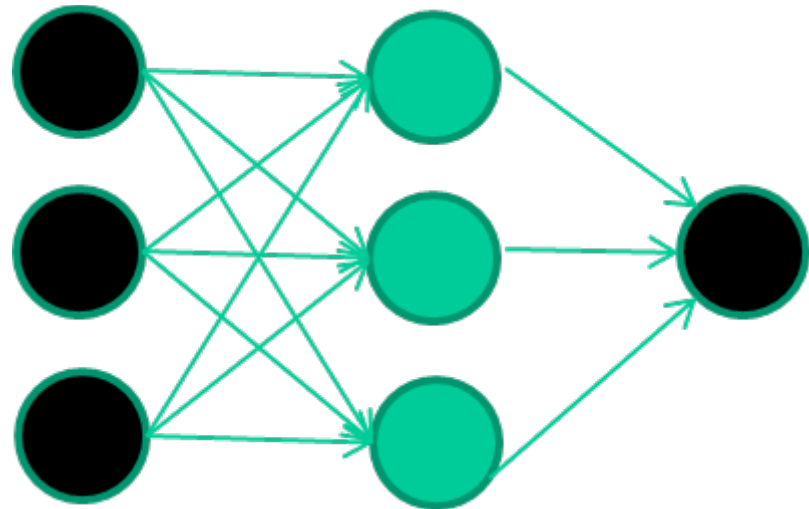
Die Ausgabefunktion

- Die Ausgabefunktion berechnet aus der Aktivierung a_k den Wert o_k , der als Ausgabe an die nächste Schicht weitergegeben wird.
- In den meisten Fällen ist die Ausgabefunktion die **Identität**, d.h. $o_k = a_k$.
- Für binäre Ausgaben wird manchmal auch eine Schwellenwertfunktion verwendet:

$$o_k = f_{out}(a_k) = \begin{cases} 1, & \text{falls } a_k \geq \theta_k, \\ 0, & \text{sonst.} \end{cases}$$

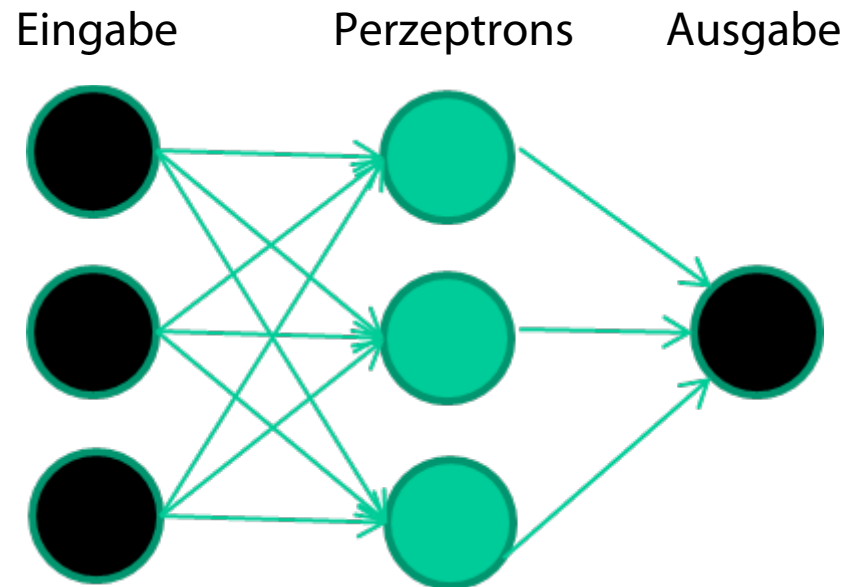
Einstellen der Gewichte mit Trainingsdaten...

<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	



Anlernen des Netzwerks

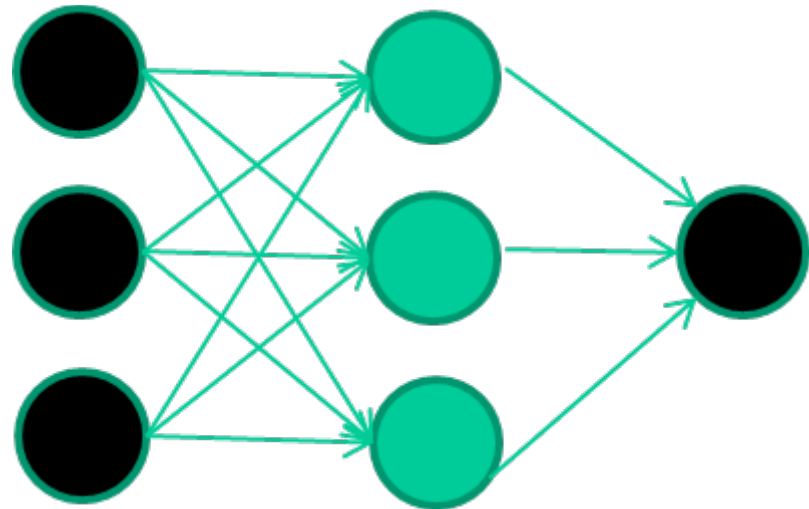
<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	



Trainingsdaten

<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	

Initialisierung mit zufälligen Gewichten



Trainingsdaten

Fields *class*

1.4 2.7 1.9 0

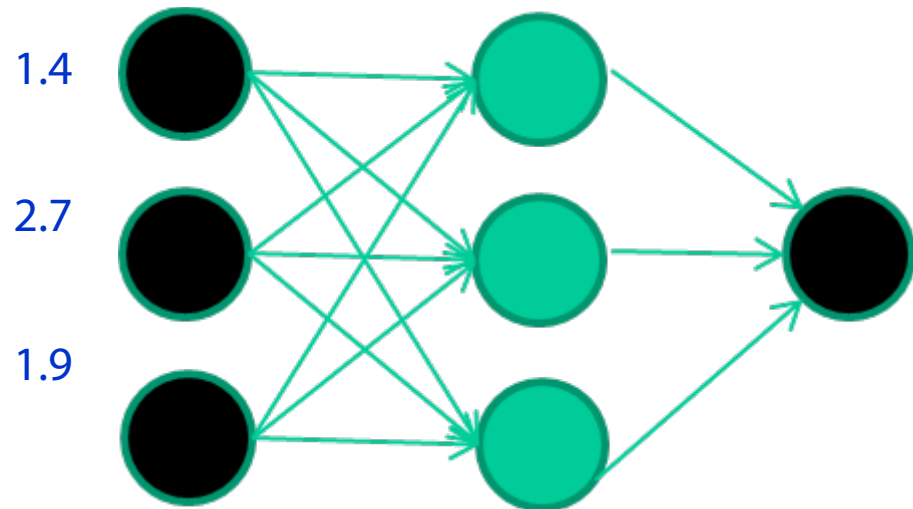
3.8 3.4 3.2 0

6.4 2.8 1.7 1

4.1 0.1 0.2 0

etc ...

Präsentierung eines Trainingsdatensatzes



Trainingsdaten

Fields *class*

1.4 2.7 1.9 0

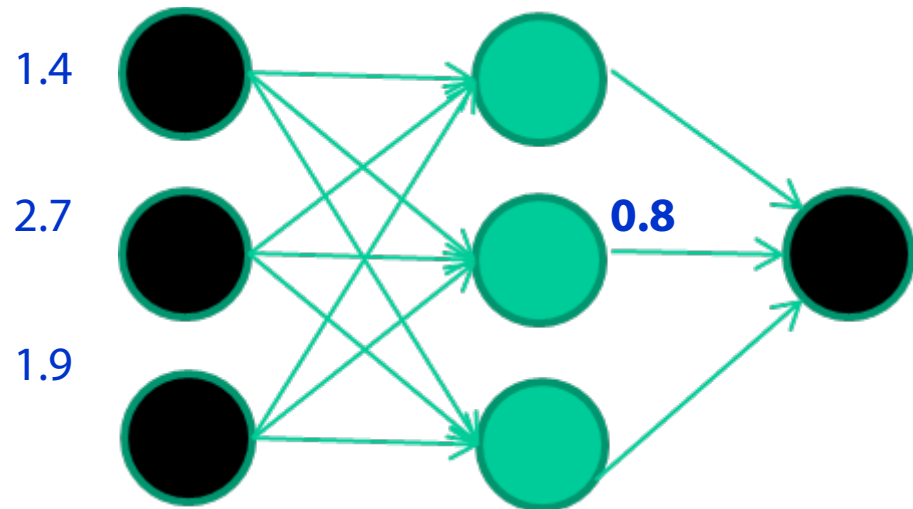
3.8 3.4 3.2 0

6.4 2.8 1.7 1

4.1 0.1 0.2 0

etc ...

Durchpropagierung zur Ausgabe



Trainingsdaten

Fields *class*

1.4 2.7 1.9 0

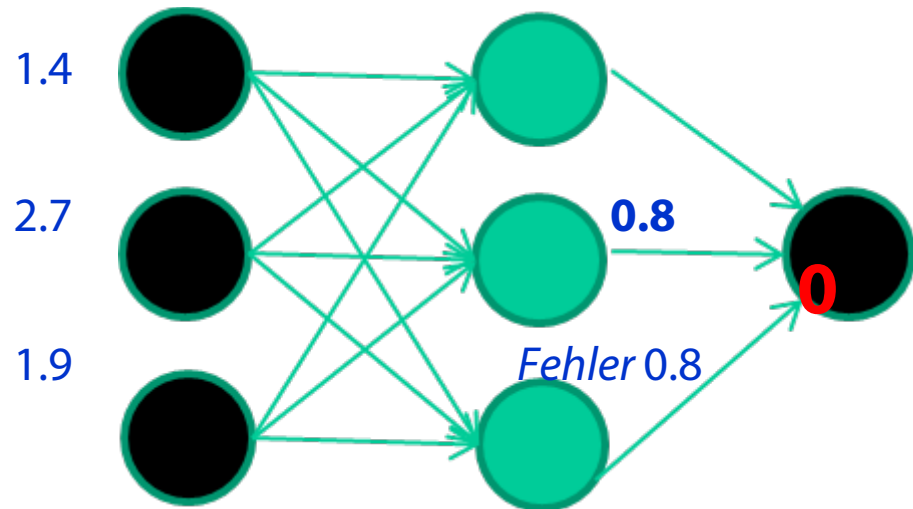
3.8 3.4 3.2 0

6.4 2.8 1.7 1

4.1 0.1 0.2 0

etc ...

Vergleich mit der Zielausgabe



Trainingsdaten

Fields *class*

1.4 2.7 1.9 0

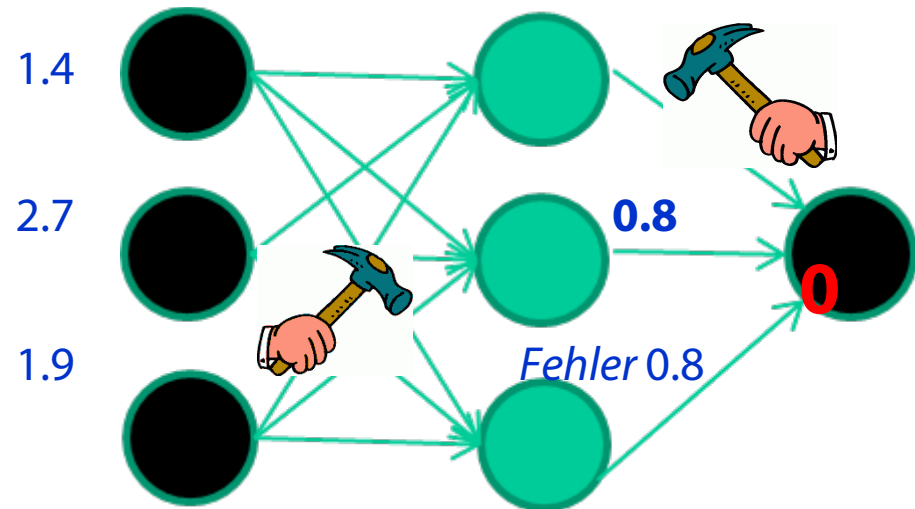
3.8 3.4 3.2 0

6.4 2.8 1.7 1

4.1 0.1 0.2 0

etc ...

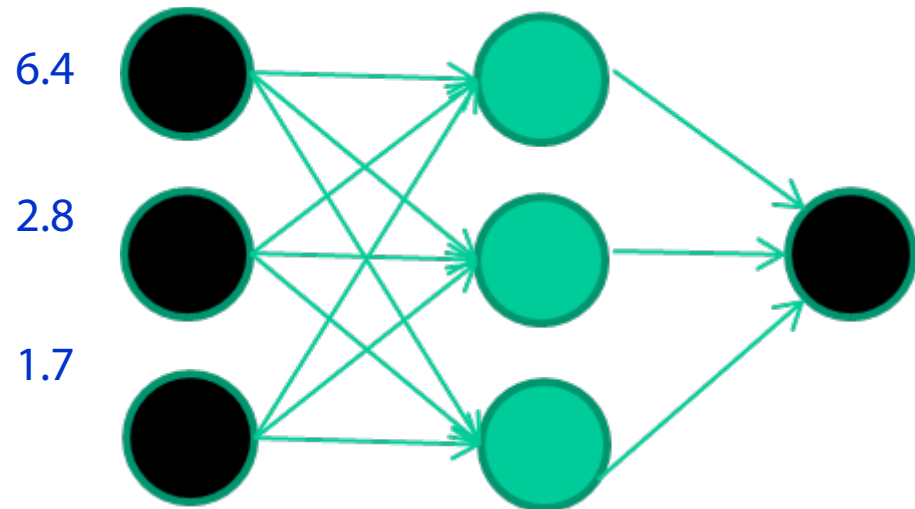
Anpassen der Gewichte gemäß Fehler



Trainingsdaten

<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	

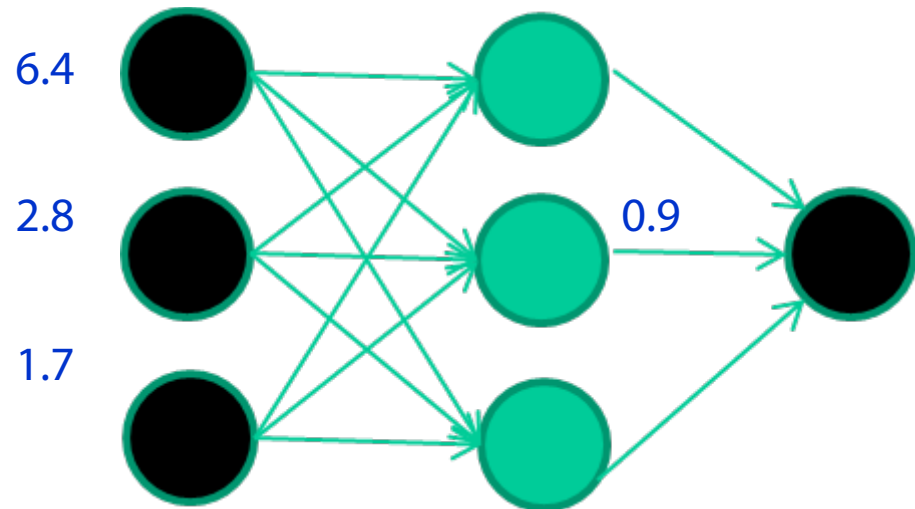
Präsentierung eines Trainingsdatensatzes



Trainingsdaten

<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	

Durchpropagierung zur Ausgabe



Trainingsdaten

Fields *class*

1.4 2.7 1.9 0

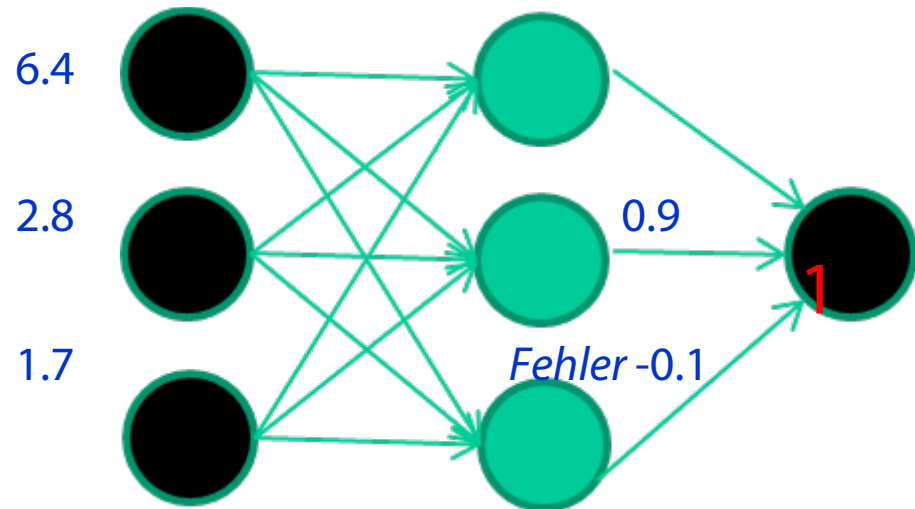
3.8 3.4 3.2 0

6.4 2.8 1.7 1

4.1 0.1 0.2 0

etc ...

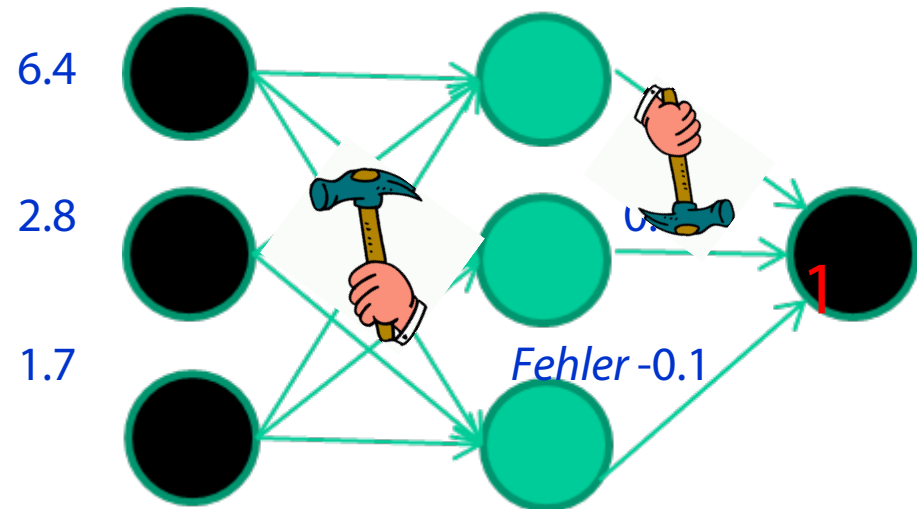
Vergleich mit der Zielausgabe



Trainingsdaten

<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	

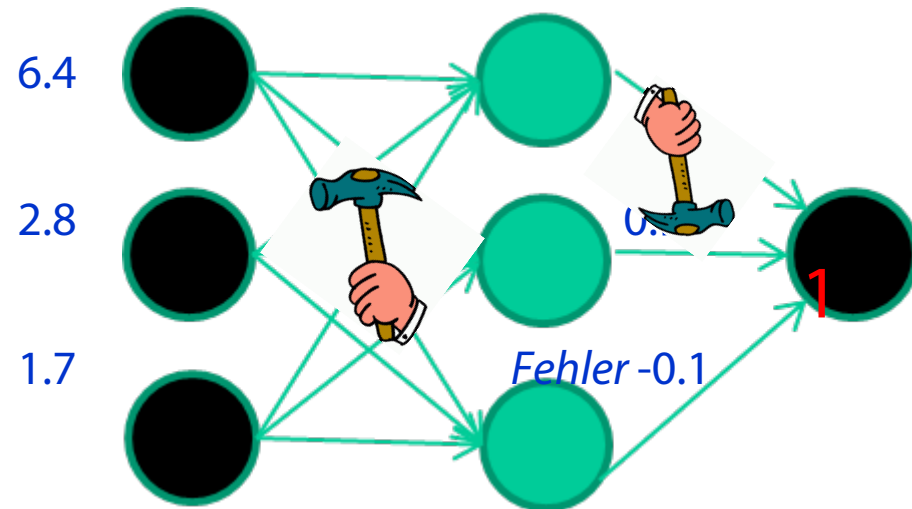
Anpassen der Gewichte gemäß Fehler



Trainingsdaten

<i>Fields</i>	<i>class</i>
1.4 2.7 1.9	0
3.8 3.4 3.2	0
6.4 2.8 1.7	1
4.1 0.1 0.2	0
etc ...	

Und so weiter



Wiederhole tausend-, vielleicht millionenmal – jedesmal mit einer zufälligen Trainingsinstanz und einer kleinen Anpassung der Gewichte

Verfahren zur Gewichts Anpassung müssen Fehler minimieren

Eine Ebene: Perzeptron-Lernregel (Delta-Lernregel)

$$w_i \leftarrow w_i + \Delta w_i$$

wobei

$$\Delta w_i \leftarrow \eta(t - o)x_i$$

und

- $t = c(\vec{x})$ der Zielwert ist
- o ist die Ausgabe des Perzeptrons
- η ist eine kleine Konstante (z.B. 0,1): die Lernrate

Gewichte häufig nur nach Verarbeitung eines ganzen Datensatzes D angepasst

Begründung für die Delta-Regel

- Idee: minimiere den quadratischen Fehler
 - D Trainingsmenge
 - t_d Wert für $d \in D$
 - o_d Ausgabe für d

$$E[\vec{w}] \equiv \frac{1}{2} \sum_{d \in D} (t_d - o_d)^2$$

- Minimumbestimmung mit 1. Ableitung

Absteigender Gradient

- Gradient

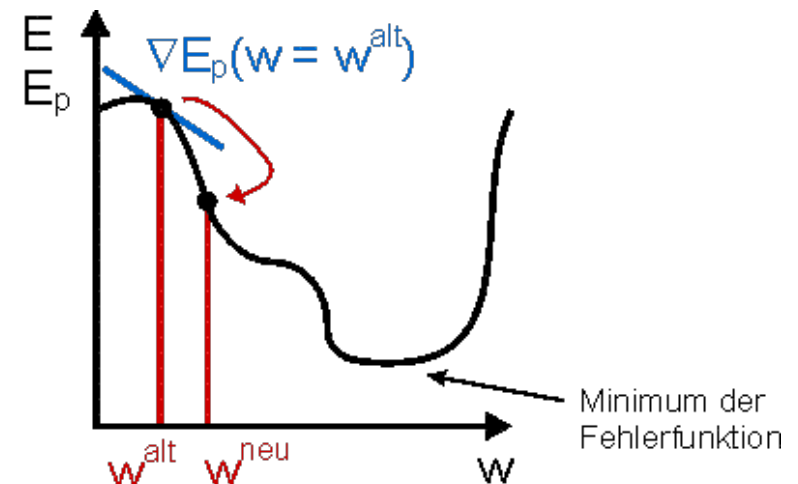
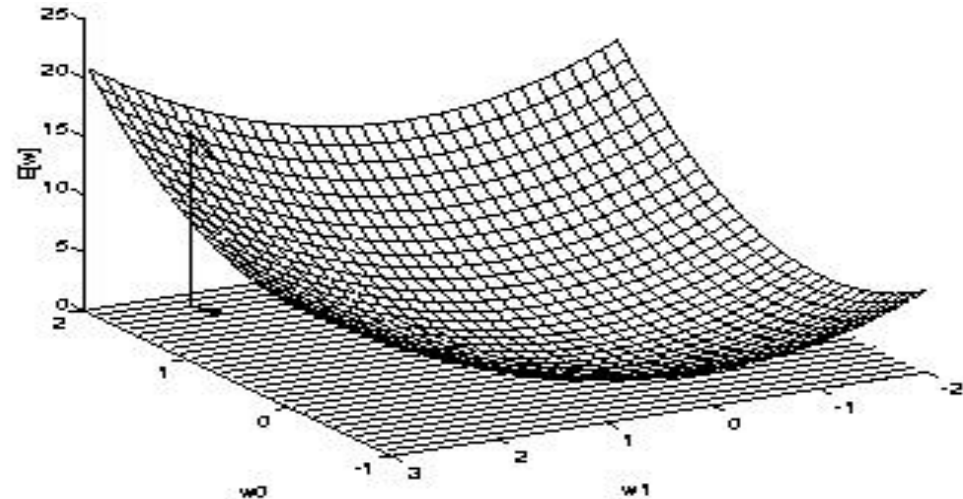
$$\nabla E[\vec{w}] \equiv \left[\frac{\partial E}{\partial w_0}, \frac{\partial E}{\partial w_1}, \dots, \frac{\partial E}{\partial w_n} \right]$$

- Lernregel

$$\Delta \vec{w} = -\eta \nabla E[\vec{w}]$$

- i.e.

$$\Delta w_i = -\eta \frac{\partial E}{\partial w_i}$$



Gradient

$$\begin{aligned}\frac{\partial E}{\partial w_i} &= \frac{\partial}{\partial w_i} \frac{1}{2} \sum_{d \in D} (t_d - o_d)^2 \\&= \frac{1}{2} \sum_{d \in D} \frac{\partial}{\partial w_i} (t_d - o_d)^2 \\&= \frac{1}{2} \sum_{d \in D} 2(t_d - o_d) \frac{\partial}{\partial w_i} (t_d - o_d) \\&= \sum_{d \in D} (t_d - o_d) \frac{\partial}{\partial w_i} (t_d - \vec{w} \cdot \vec{x}_d) \\&= \sum_{d \in D} (t_d - o_d) (-x_{i,d}) = - \sum_{d \in D} (t_d - o_d) x_{i,d}\end{aligned}$$

Absteigender Gradient (Forts.)

- Gradient

$$\nabla E[\vec{w}] \equiv \left[\frac{\partial E}{\partial w_0}, \frac{\partial E}{\partial w_1}, \dots, \frac{\partial E}{\partial w_n} \right]$$

$$\frac{\partial E}{\partial w_i} = - \sum_{d \in D} (t_d - o_d) x_{i,d}$$

- Lernregel

$$\Delta \vec{w} = -\eta \nabla E[\vec{w}]$$

- i.e.

$$\Delta w_i = -\eta \frac{\partial E}{\partial w_i}$$

Iterativ pro Datenpunkt d

$$\Delta w_i = \eta (t - o) x_i$$

$$\Delta w_i = \eta \sum_{d \in D} (t_d - o_d) x_{i,d}$$

Algorithmus für *eine* Gewichts Anpassung

- Jedes Trainingsbeispiel sei ein Paar $\langle \vec{x}, t \rangle$
 - $\vec{x}$ ist ein Inputvektor (Vektor mit Feature-Werten)
 - t ist der Zielwert (Klassenlabel/Ausgabewert)
 - η ist die Lernrate
- Initialisiere jedes w_i mit einem beliebigen, kleinen Wert
- Bis die Abbruchbedingung erfüllt ist (z.B. Anzahl an verarbeiteten Trainingsbeispielen):
 - Initialisiere jedes Δw_i mit 0
 - Für jedes Trainingsbeispiel $\langle \vec{x}, t \rangle$
 - Berechne o_t
 - Für jedes Gewicht w_i : $\Delta w_i \leftarrow \Delta w_i + \eta(t - o_t)x_i$
 - Für jedes w_i : $w_i \leftarrow w_i + \Delta w_i$

Perzeptron-Lernregel

Man kann zeigen, dass der Vorgang konvergiert, ...

- ... wenn die Daten linear separierbar sind
- ... und η genügend klein gewählt wird

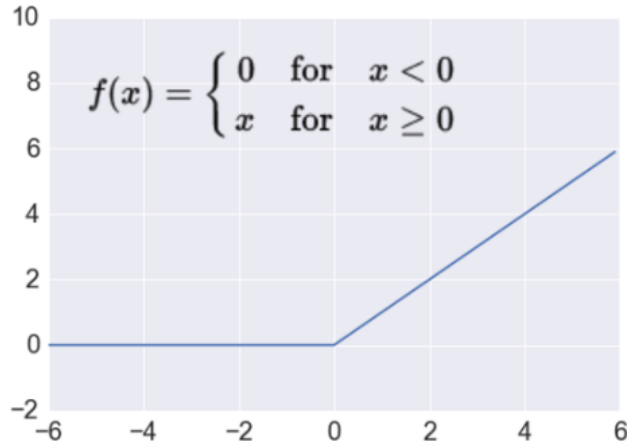
Schon früher untersucht:

D. Hebb: The organization of behavior. A neuropsychological theory.
Erlbaum Books, Mahwah, N.J., 1949

Netzwerke daher von manchen
als künstliche neuronale Netze bezeichnet

Später für mehrschichtige Netze erweitert (Deep Learning):
Fehlerrückführung durch mehrere Ebenen (Backpropagation)

Aktivierung: ReLU



<http://adilmoujahid.com/images/activation.png>

Takes a real-valued number and thresholds it at zero $f(x) = \max(0, x)$

$$R^n \rightarrow R_+^n$$

Most Deep Networks use ReLU nowadays

Trains much **faster**

- accelerates the convergence of SGD
- due to linear, non-saturating form

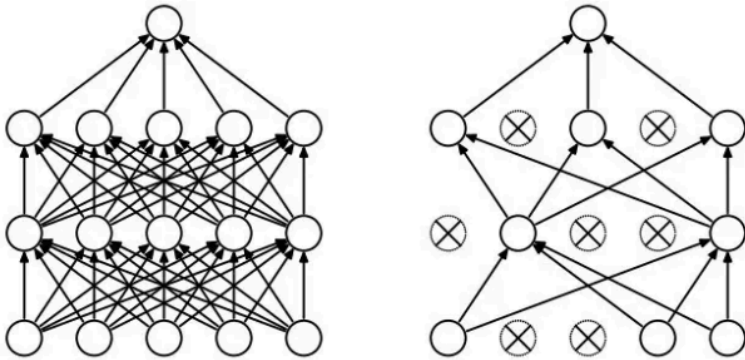
Less expensive operations

- compared to sigmoid/tanh (exponentials etc.)
- implemented by simply thresholding a matrix at zero

More **expressive**

Prevents the **gradient vanishing problem**

Regularization



Dropout

- Randomly drop units (along with their connections) during training
- Each unit retained with fixed probability p , independent of other units
- **Hyper-parameter** p to be chosen (tuned)

Srivastava, Nitish, et al. ["Dropout: a simple way to prevent neural networks from overfitting."](#) Journal of machine learning research (2014)

L2 = weight decay

- Regularization term that penalizes big weights, added to the objective
- Weight decay value determines how dominant regularization is during gradient computation
- Big weight decay coefficient $\rightarrow$ big penalty for big weights

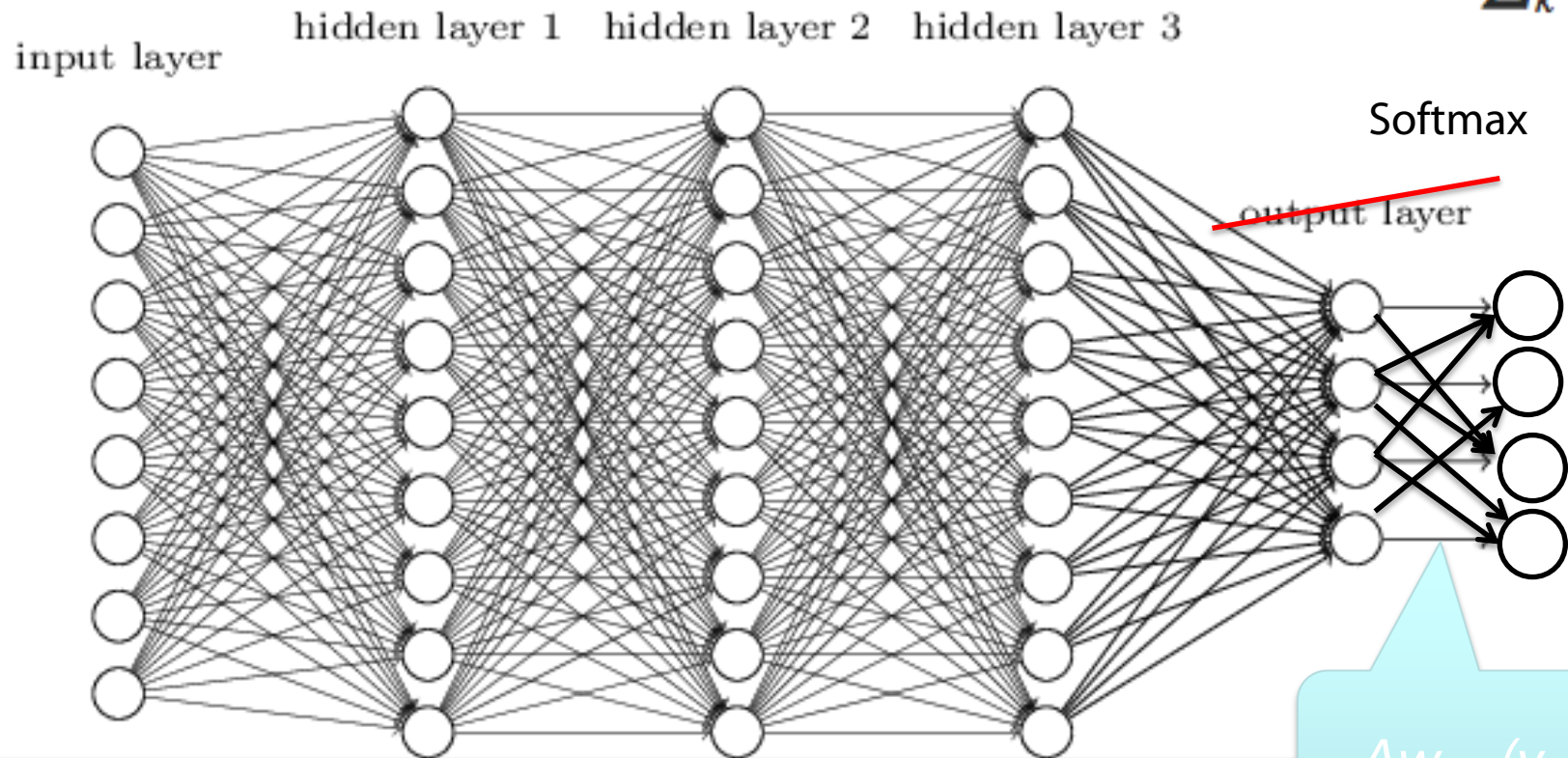
$$J_{reg}(\theta) = J(\theta) + \lambda \sum_k \theta_k^2$$

Early-stopping

- Use validation error to decide when to stop training
- Stop when monitored quantity has not improved after n subsequent epochs
- n is called patience

Softmax output layer

$$a_j^L = \frac{e^{z_j^L}}{\sum_k e^{z_k^L}},$$



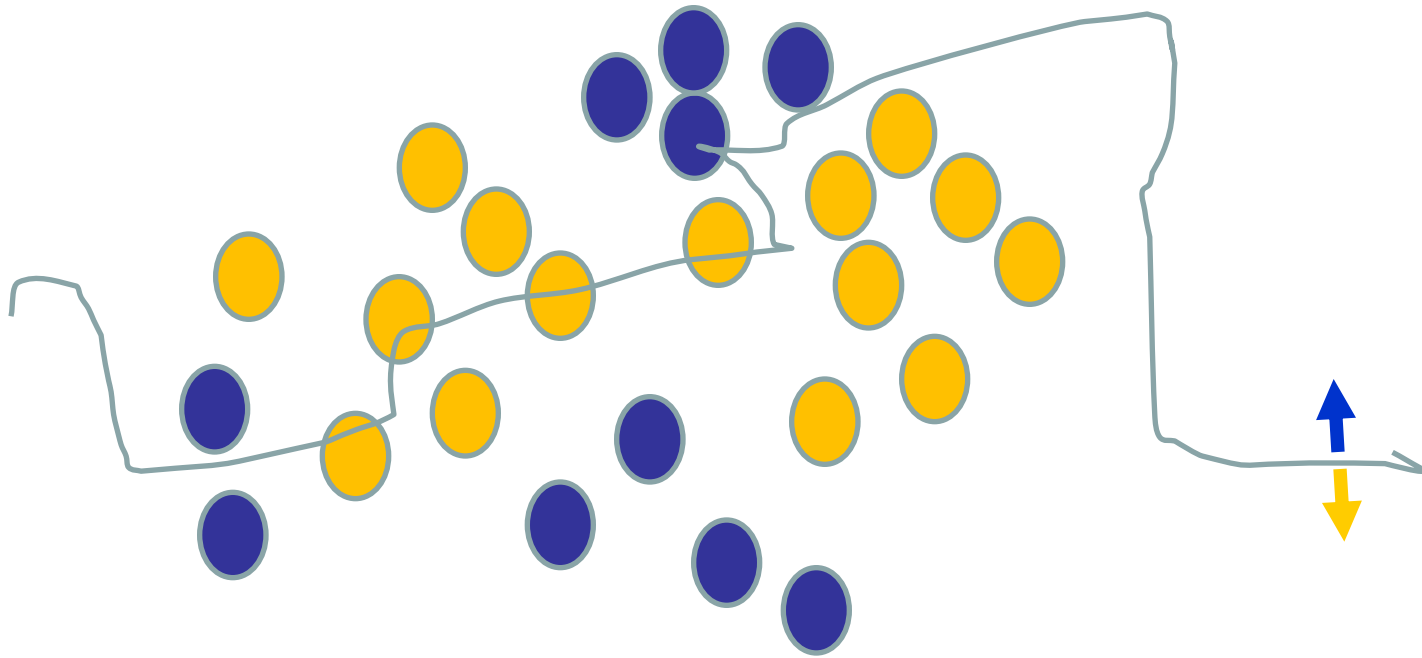
Network outputs a probability distribution!

Cross-entropy loss after a softmax layer gives a very simple, numerically stable gradient

$$\Delta w_{ij} = (y_i - z_i) y_j$$

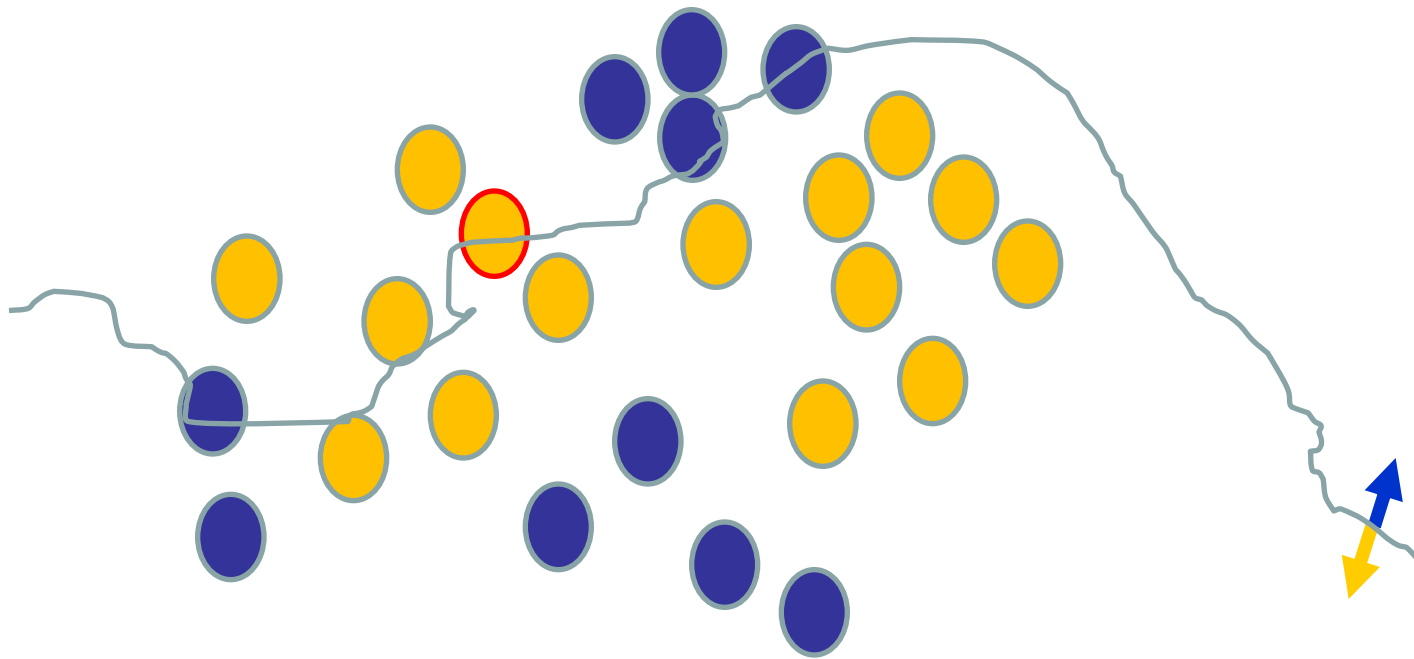
Perspektive der Entscheidungsgrenzenanpassung

Zufällige Initialgewichte



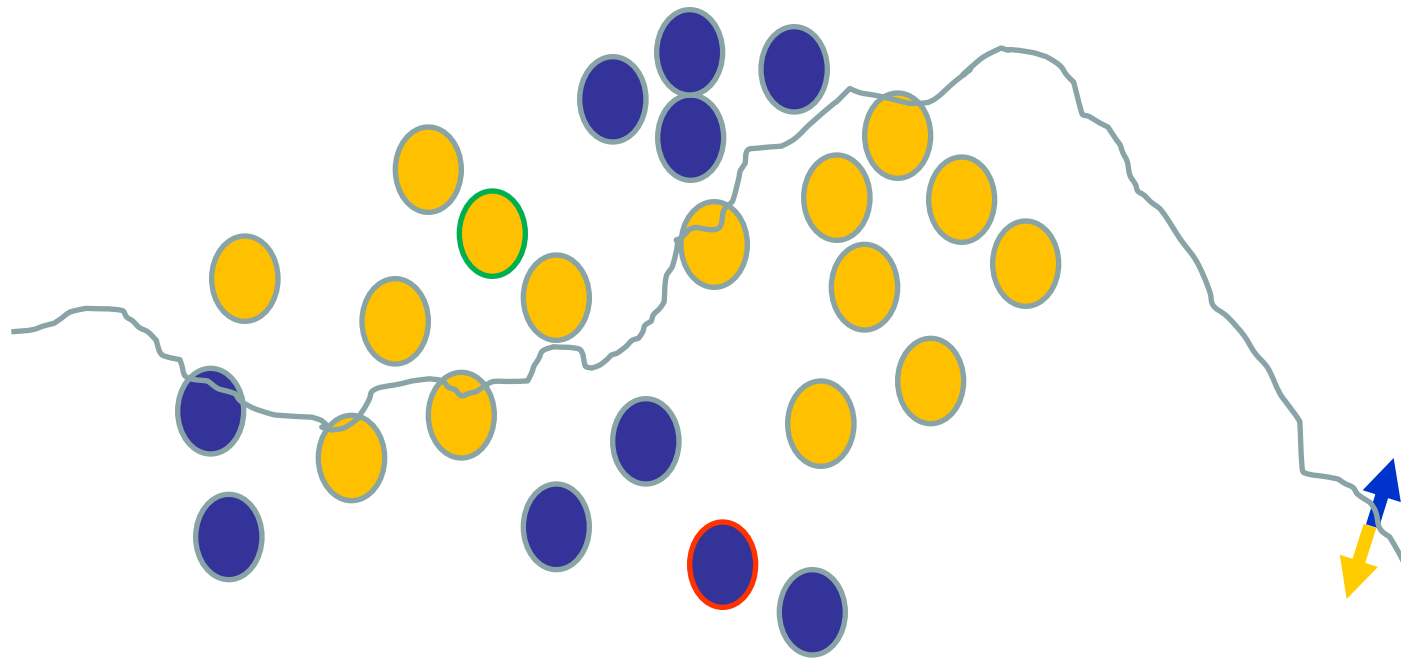
Perspektive der Entscheidungsgrenzenanpassung

Verwenden einer Trainingsinstanz / Anpassung der Gewichte



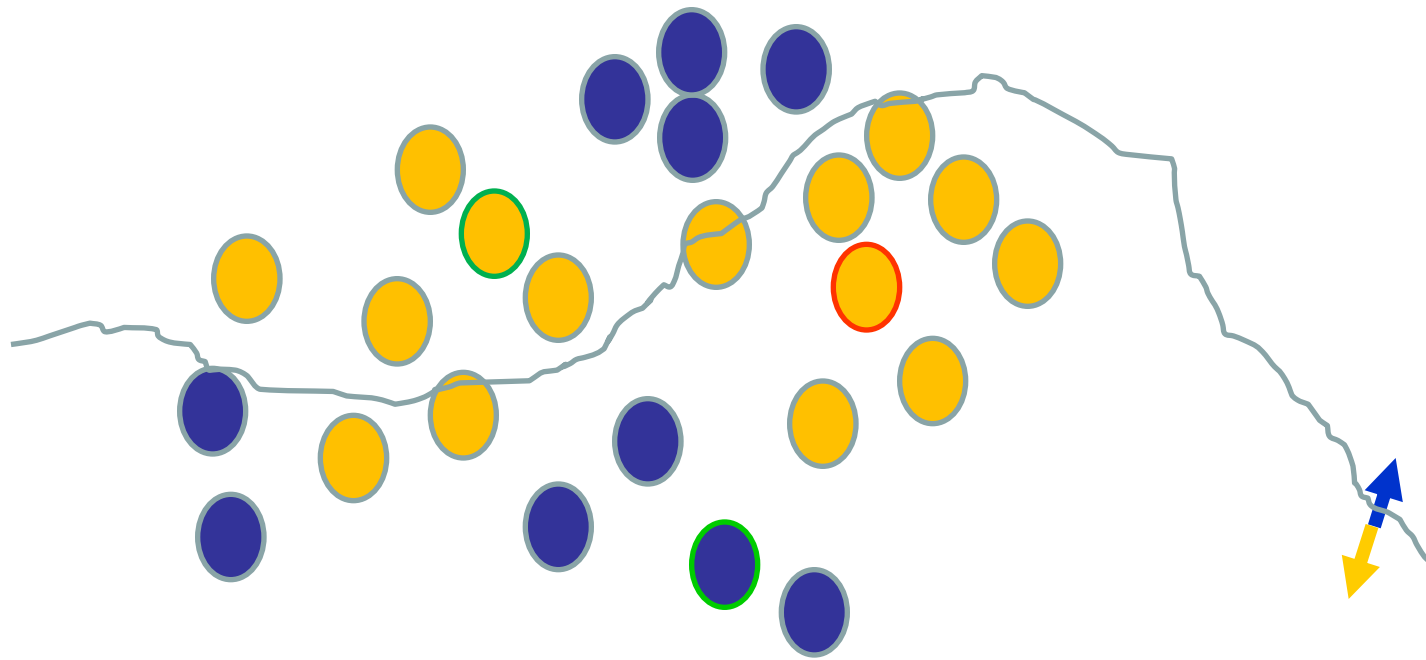
Perspektive der Entscheidungsgrenzenanpassung

Verwenden einer Trainingsinstanz / Anpassung der Gewichte



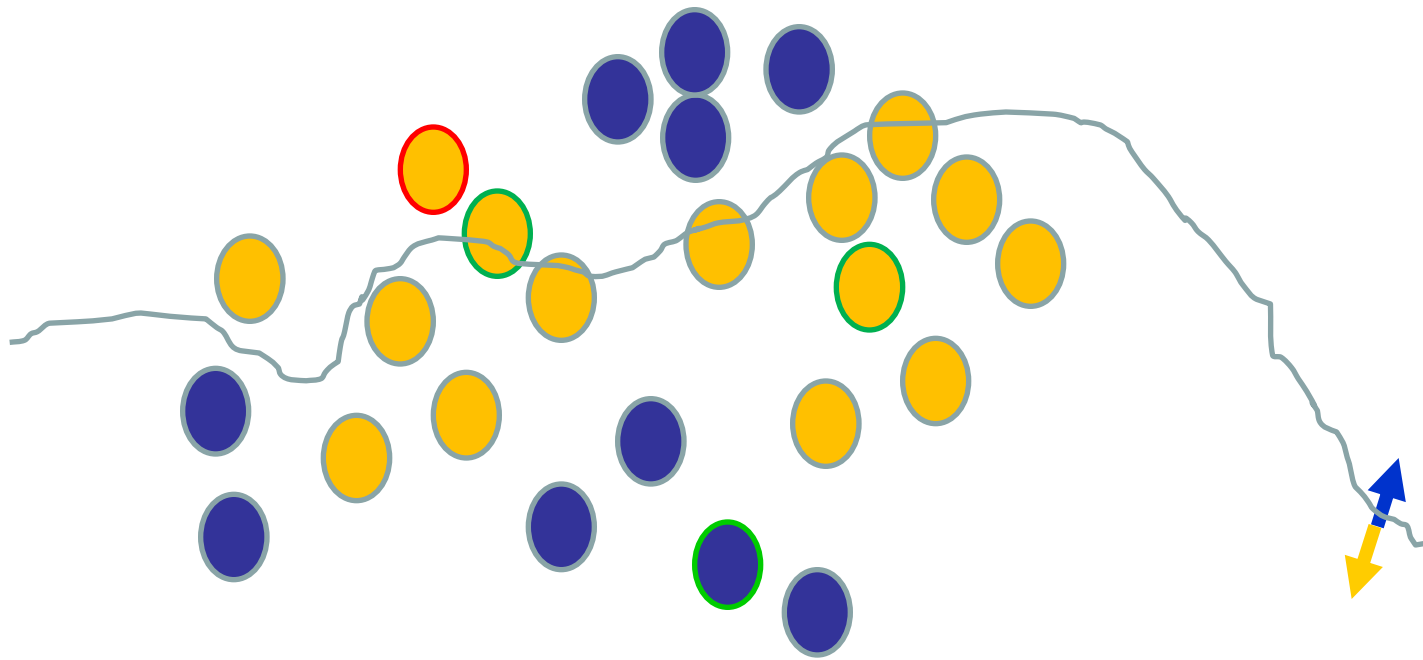
Perspektive der Entscheidungsgrenzenanpassung

Verwenden einer Trainingsinstanz / Anpassung der Gewichte



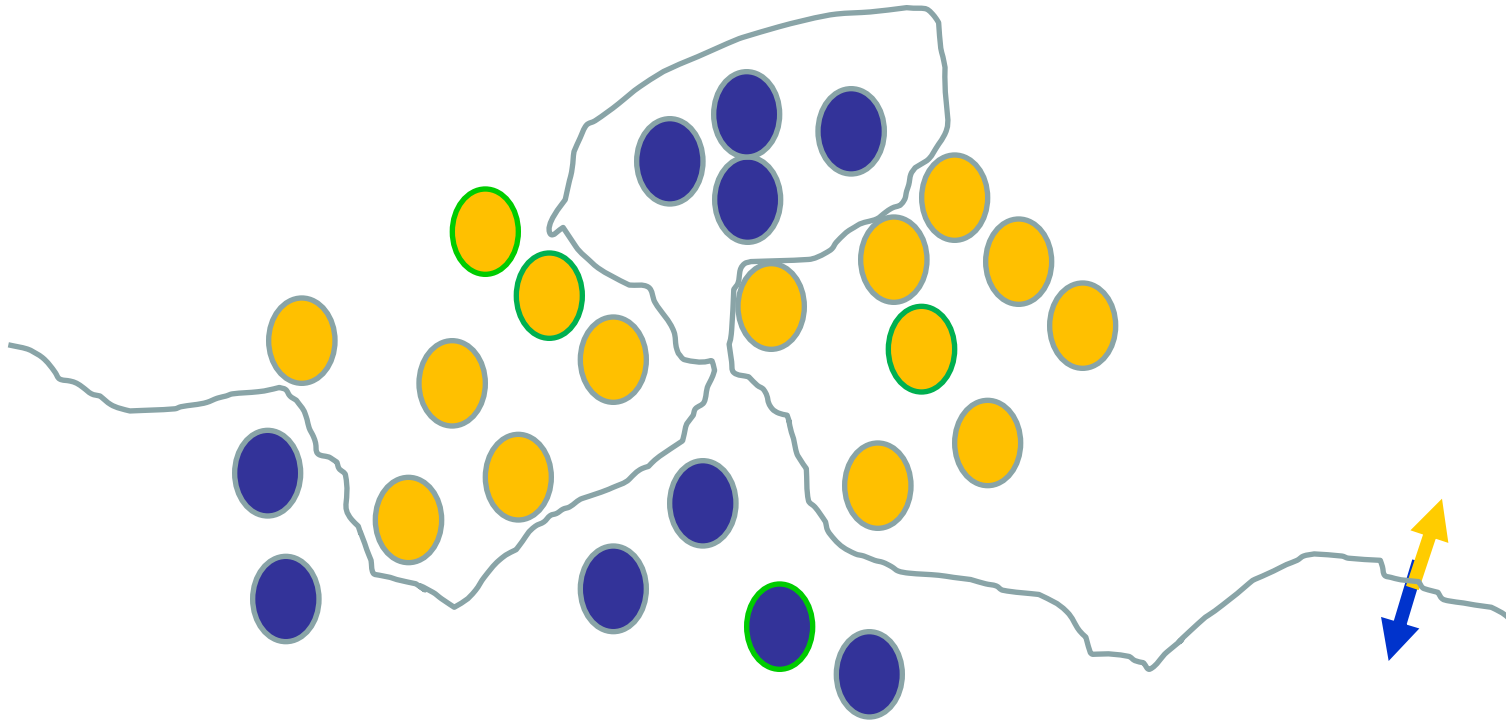
Perspektive der Entscheidungsgrenzenanpassung

Verwenden einer Trainingsinstanz / Anpassung der Gewichte

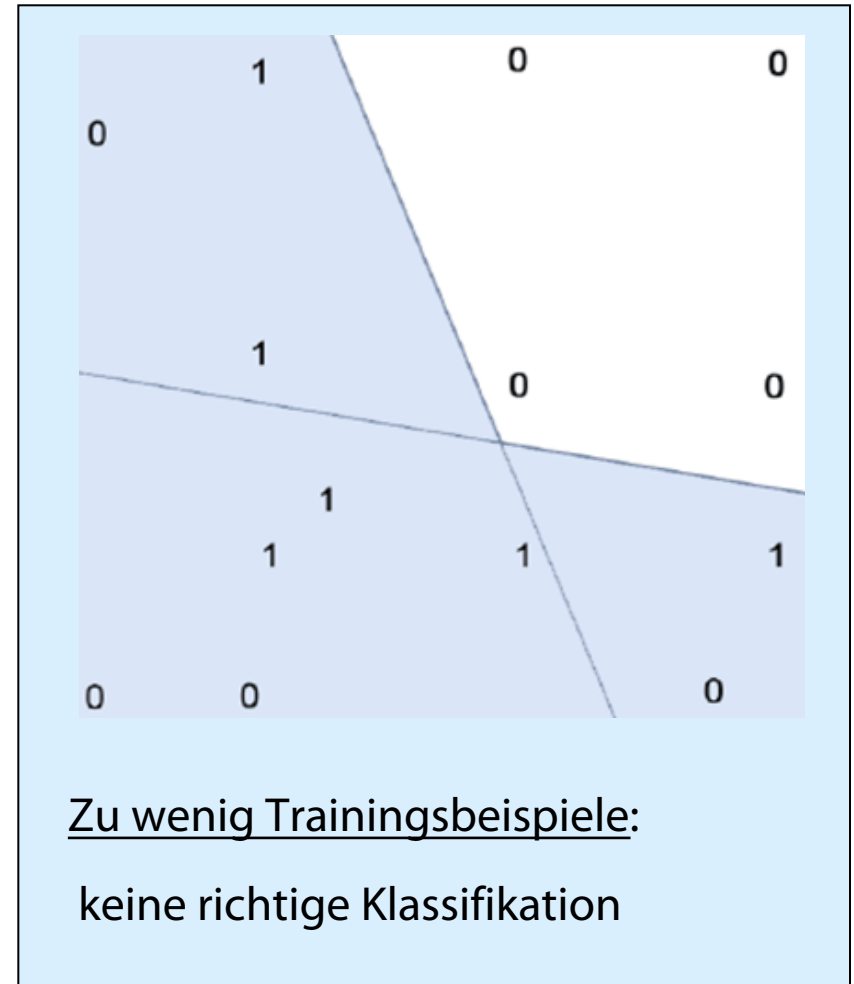
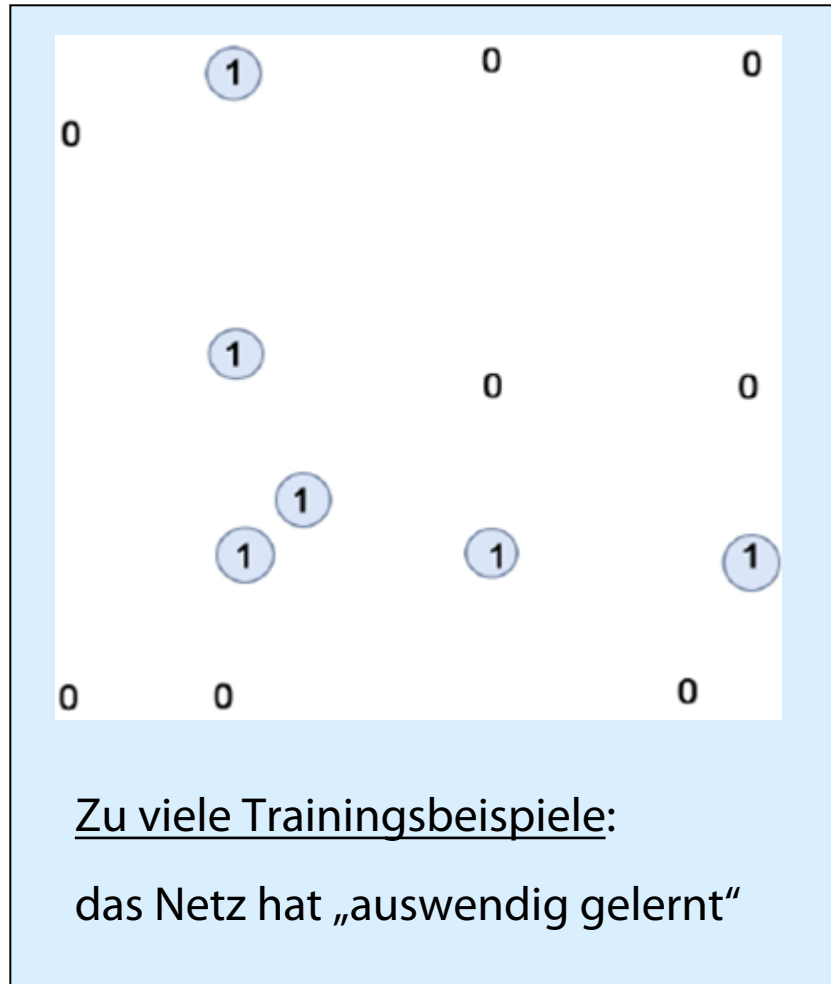


Perspektive der Entscheidungsgrenzenanpassung

Und schließlich....

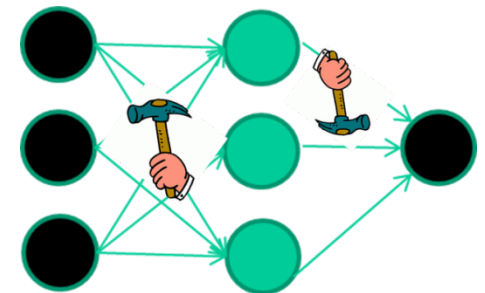


Fehler bei einer ungeeigneten Trainingsmenge



Lernverfahren: Resümee

- Epoche: einmalige Verwendung des gesamten Datensatzes
 - Manchmal zu groß
- Batch: Zerlegung eines Datensatzes in Teilmengen
 - Für jede Teilmenge: Passe Gewichte an
 - Anzahl der Teilmengen: #Iterationen
- Tausende von kleinen Anpassungen, jede macht das Netz besser für die letzte Eingabe (aber vielleicht schlechter für frühere Eingaben)
- Verwende mehrere Epochen
- Wieviele Epochen? Batches?
 - Ausprobieren
- Durch verdammt gutes Glück kommt oft eine Funktion heraus, die gut genug in Anwendungen funktioniert



Word-Word Associations in Document Retrieval

Recap bag-of-words approaches

- Client profiles, TF-IDF

Words are not independent of each other

Need to represent some aspects of word semantics

Point(wise) Mutual Information: PMI

- **Measure of association** used in information theory and statistics

$$\text{pmi}(x; y) \equiv \log \frac{p(x, y)}{p(x)p(y)} = \log \frac{p(x|y)}{p(x)} = \log \frac{p(y|x)}{p(y)}$$

- Positive PMI: $\text{PPMI}(x, y) = \max(\text{pmi}(x, y), 0)$
- Quantifies the discrepancy between the **probability of their coincidence** given their **joint distribution** and their **individual distributions**, assuming independence
- **Finding collocations and associations between words**
- **Countings** of occurrences and co-occurrences of words in a text corpus can be used to **approximate the probabilities** $p(x)$ or $p(y)$ and $p(x,y)$ respectively

PMI – Example

word 1	word 2	count word 1	count word 2	count of co-occurrences	PMI
puerto	rico	1938	1311	1159	10.0349081703
hong	kong	2438	2694	2205	9.72831972408
los	angeles	3501	2808	2791	9.56067615065
carbon	dioxide	4265	1353	1032	9.09852946116
prize	laureate	5131	1676	1210	8.85870710982
san	francisco	5237	2477	1779	8.83305176711
nobel	prize	4098	5131	2498	8.68948811416
ice	hockey	5607	3002	1933	8.6555759741
star	trek	8264	1594	1489	8.63974676575
car	driver	5578	2749	1384	8.41470768304
it	the	283891	3293296	3347	-1.72037278119
are	of	234458	1761436	1019	-2.09254205335
this	the	199882	3293296	1211	-2.38612756961
is	of	565679	1761436	1562	-2.54614706831
and	of	1375396	1761436	2949	-2.79911817902
a	and	984442	1375396	1457	-2.92239510038
in	and	1187652	1375396	1537	-3.05660070757
to	and	1025659	1375396	1286	-3.08825363041
to	in	1025659	1187652	1066	-3.12911348956
of	and	1761436	1375396	1190	-3.70663100173

- Counts of pairs of words getting the **most and the least PMI scores** in the first 50 millions of words in **Wikipedia** (dump of October 2015)
- Filtering by 1,000 or more co-occurrences.
- The frequency of each count can be obtained by dividing its value by 50,000,952. (Note: natural log is used to calculate the PMI values in this example, instead of log base 2)

PMI – Co-occurrence Matrix

	Add-2 Smoothed Count(w,context)				
	computer	data	pinch	result	sugar
apricot	2	2	3	2	3
pineapple	2	2	3	2	3
digital	4	3	2	3	2
information	3	8	2	6	2

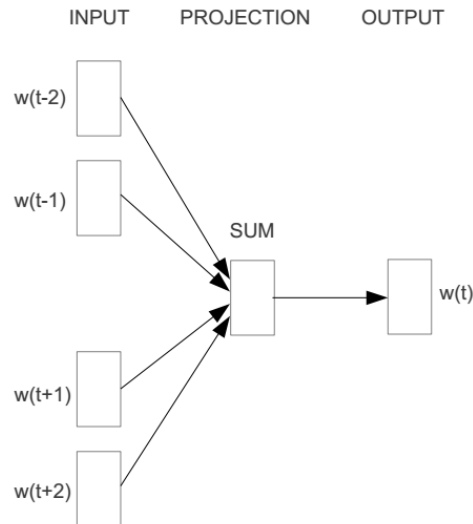
	PPMI(w,context)				
	computer	data	pinch	result	sugar
apricot	-	-	2.25	-	2.25
pineapple	-	-	2.25	-	2.25
digital	1.66	0.00	-	0.00	-
information	0.00	0.57	-	0.47	-

Embedding Approaches to Word Semantics

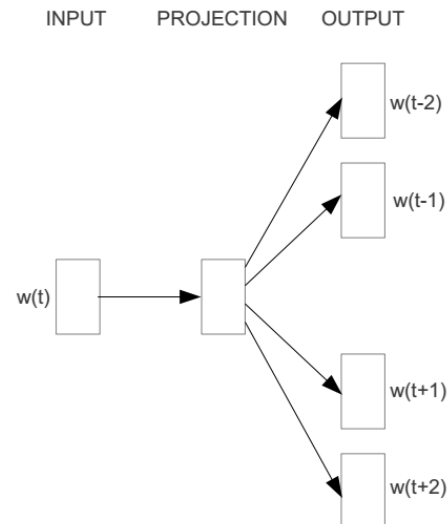
- Represent each word with a low-dimensional vector
- Word similarity = vector similarity
- Key idea: Predict surrounding words of every word

Represent the meaning of **words** – word2vec

- 2 basic structural models:
 - Continuous Bag of Words (CBOW): use a window of words to predict the middle word
 - Skip-gram (SG): use a word to predict the surrounding ones in window.



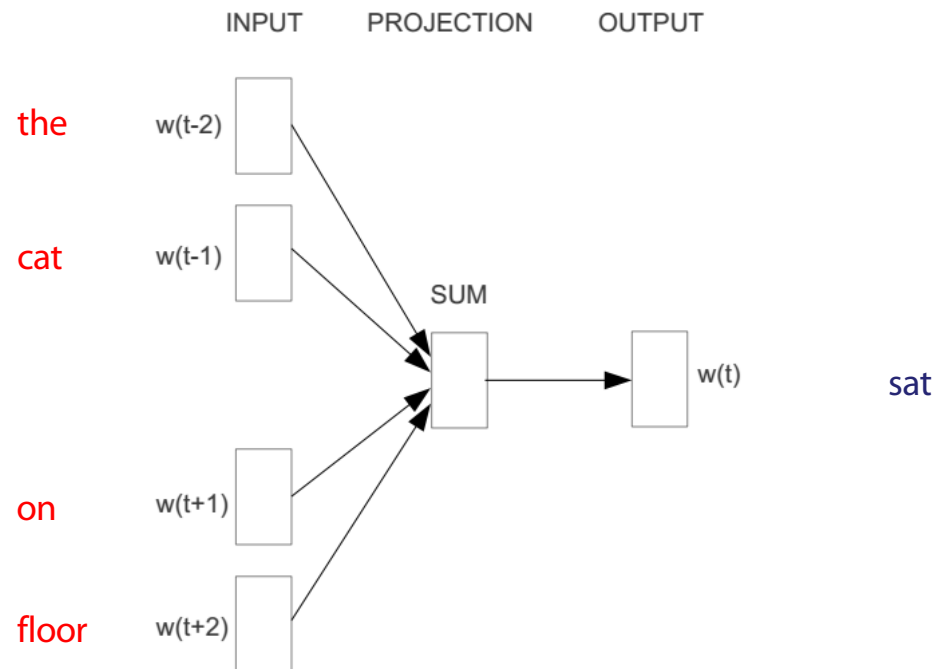
CBOW

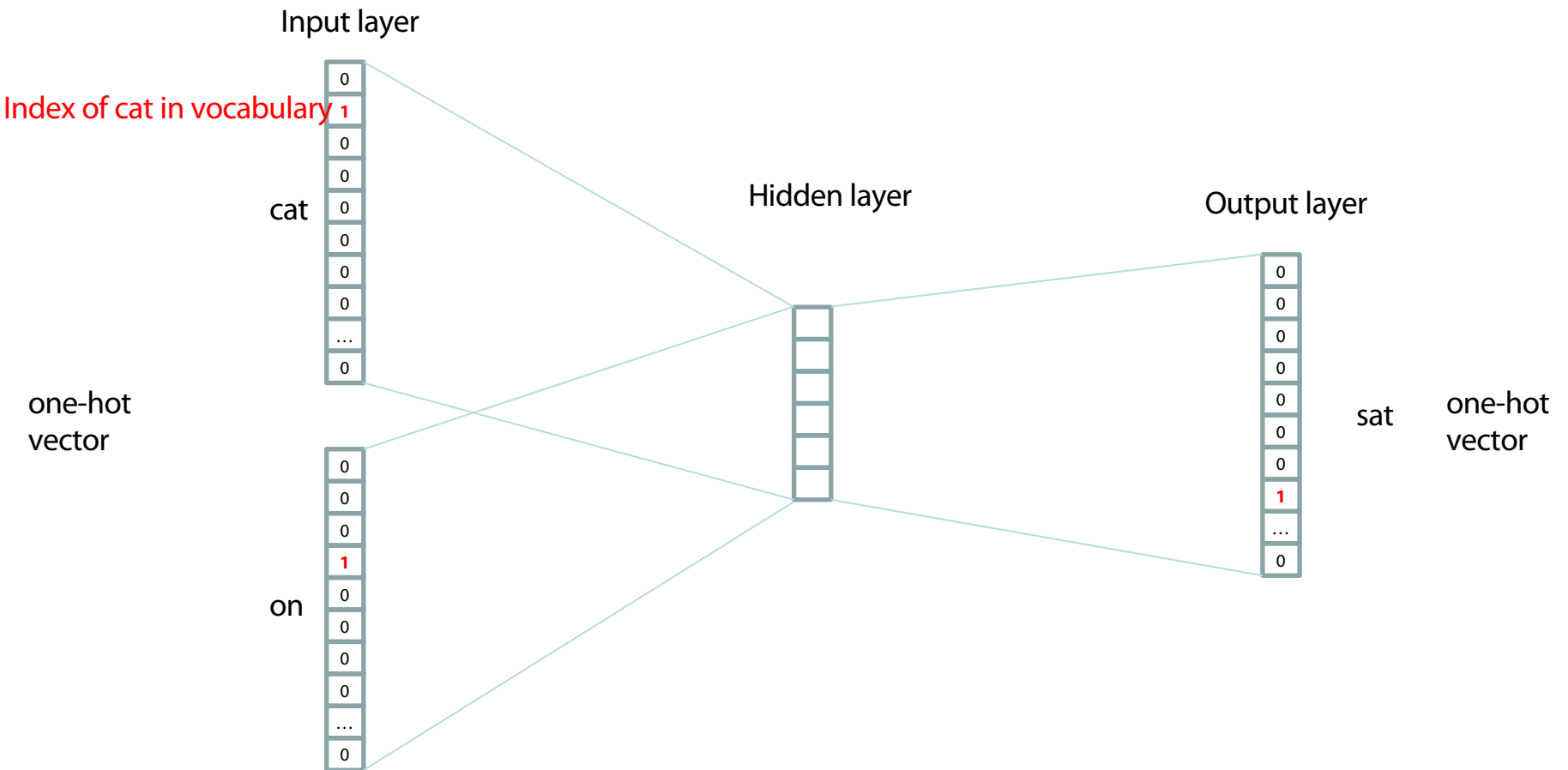


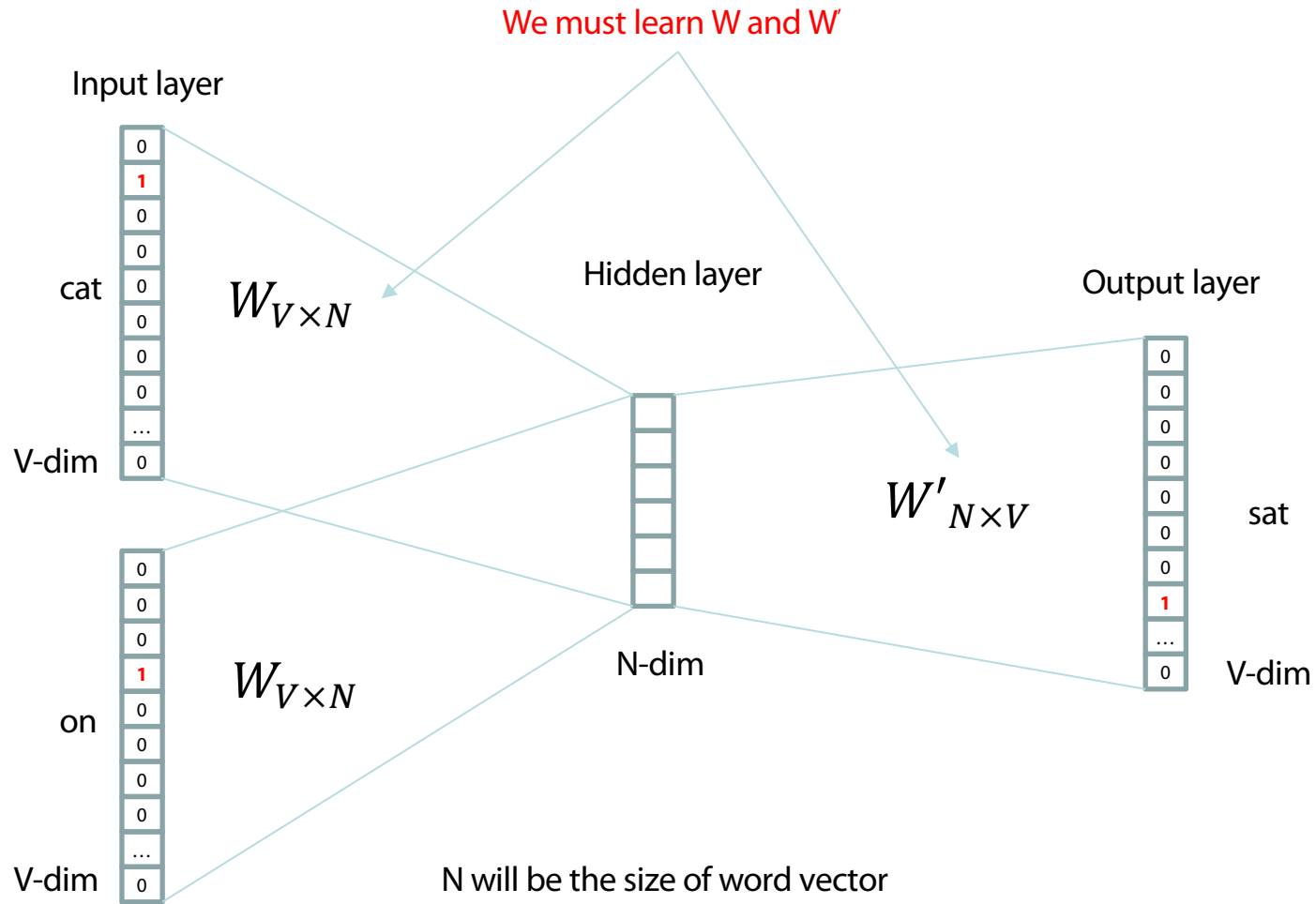
Skip-gram

Word2vec – Continuous Bag of Word

- E.g. “The cat <sat> on floor”
 - Window size = 2

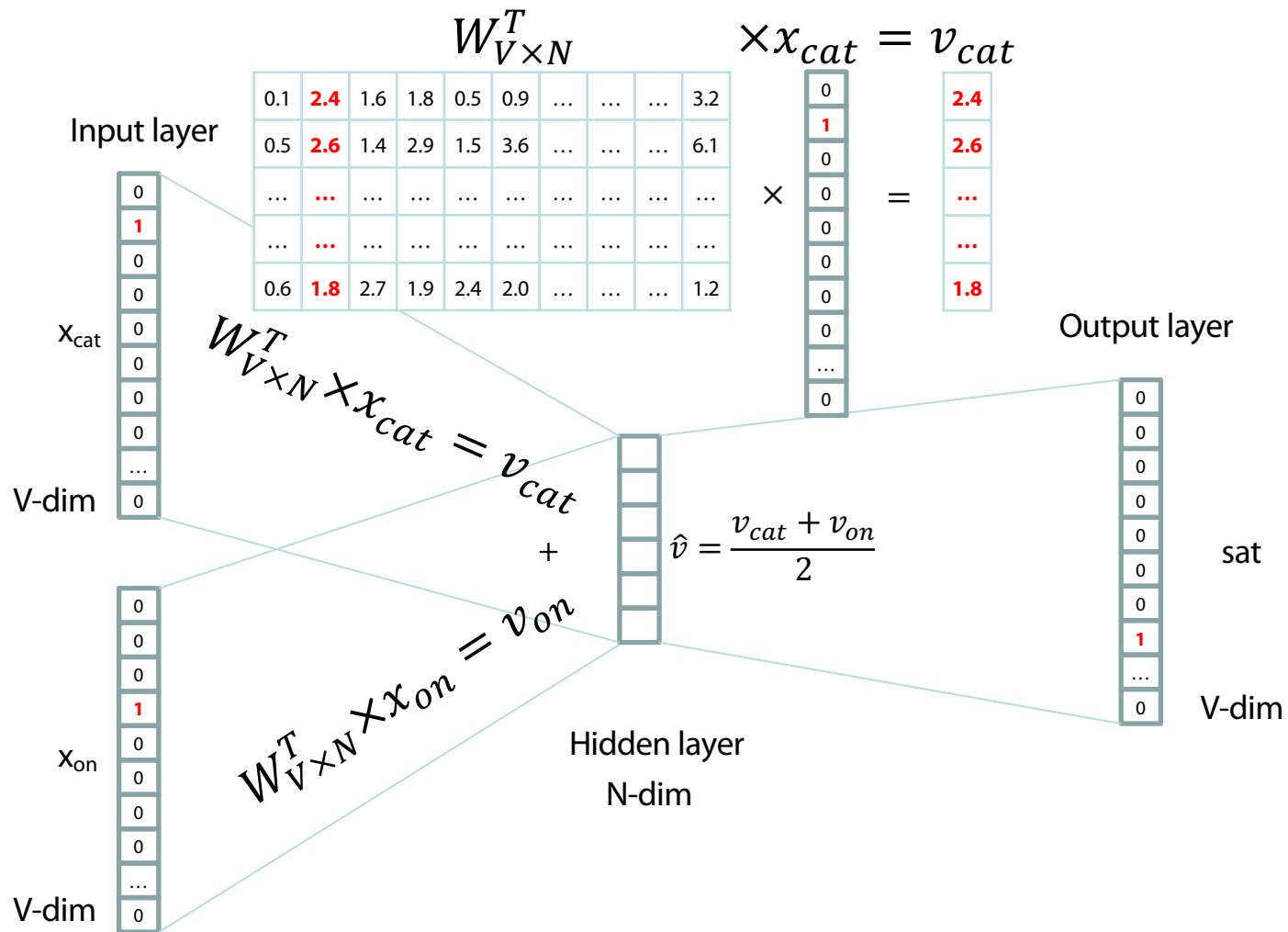


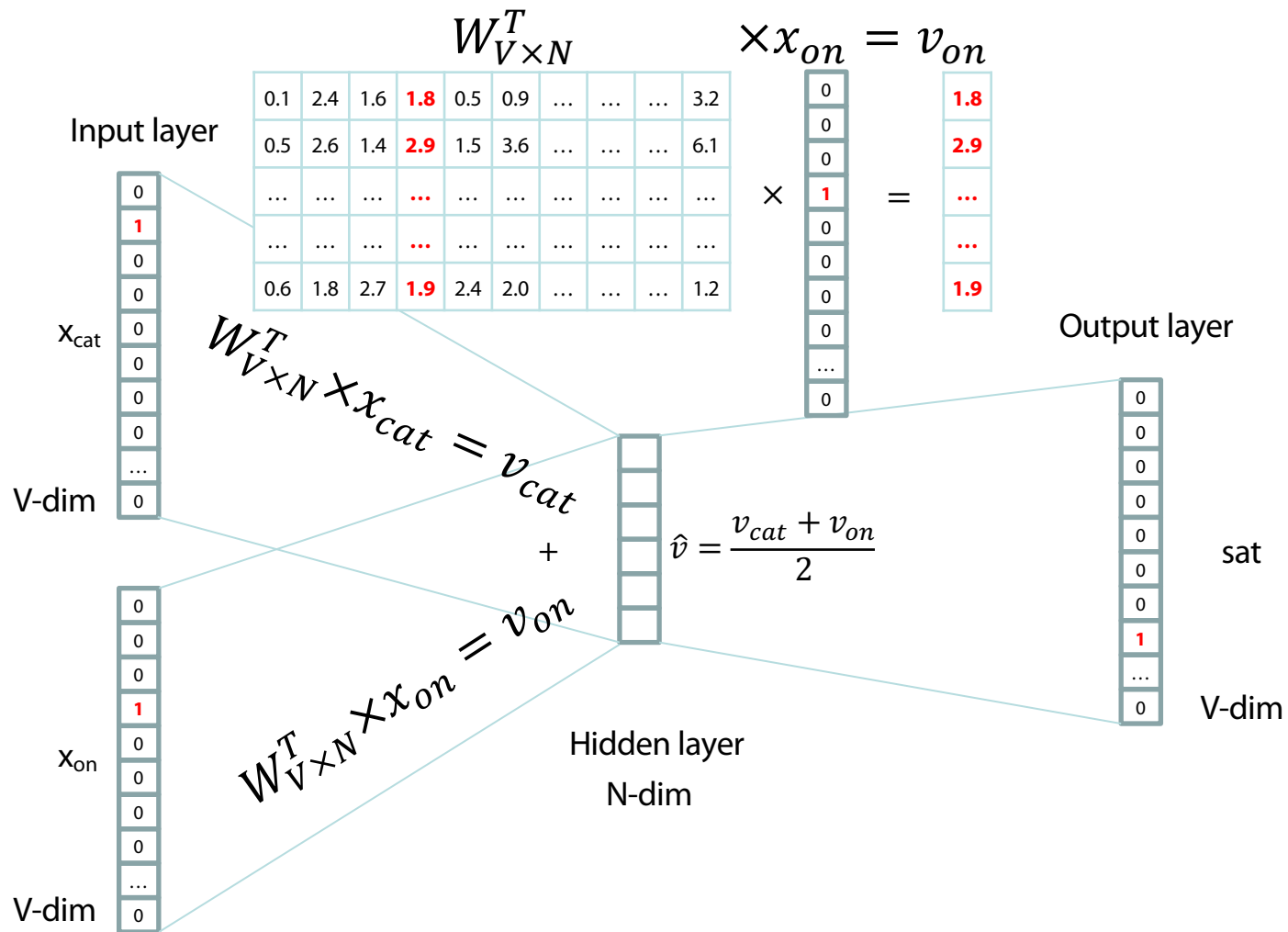


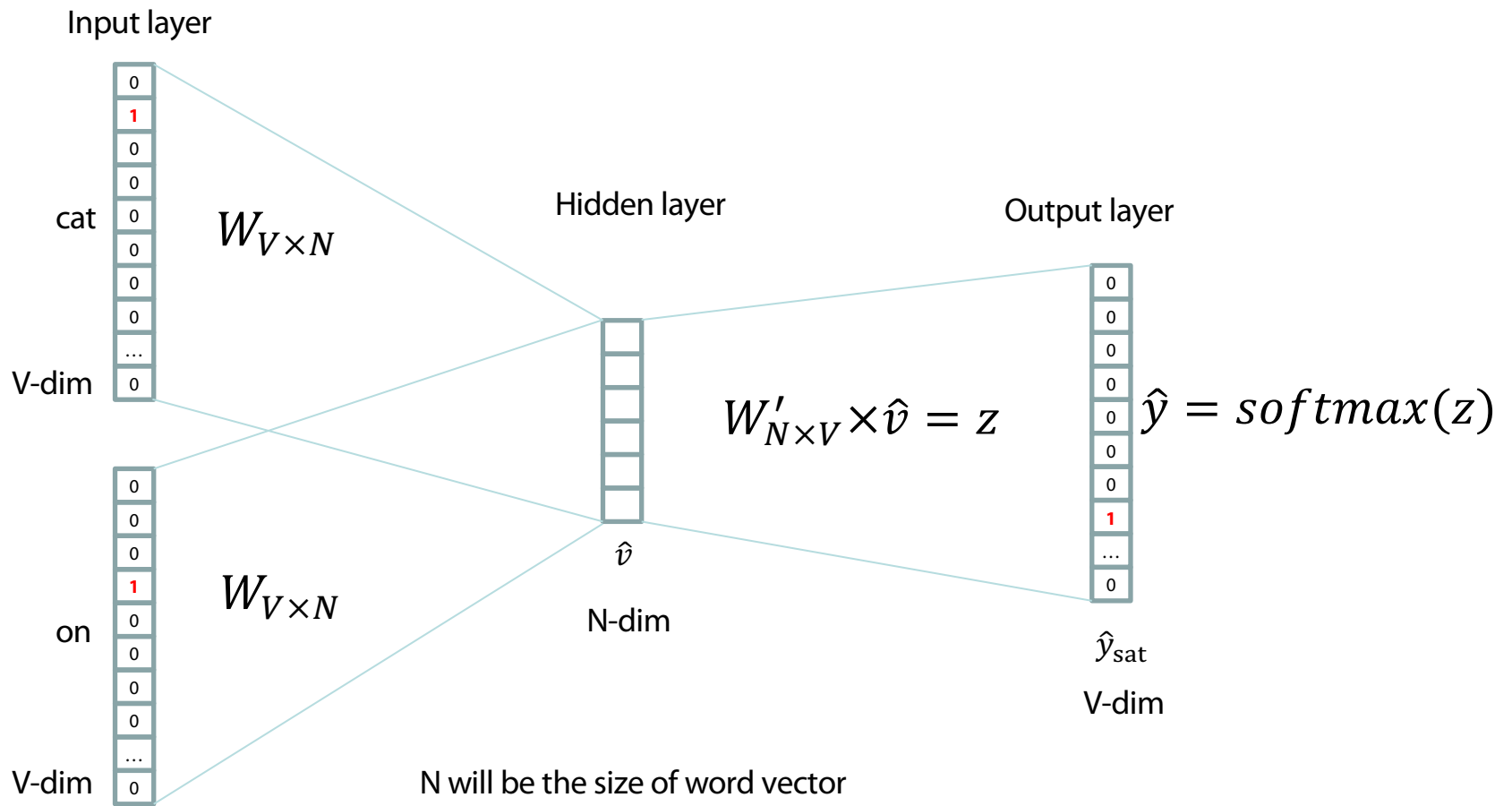


Deep Learning

- Hidden layer represents feature space
 - Making explicit features in the data...
 - ... that are relevant for a certain task
- Determine features automatically
 - Learning suitable mappings into feature space
- Deep learning also known as representation learning







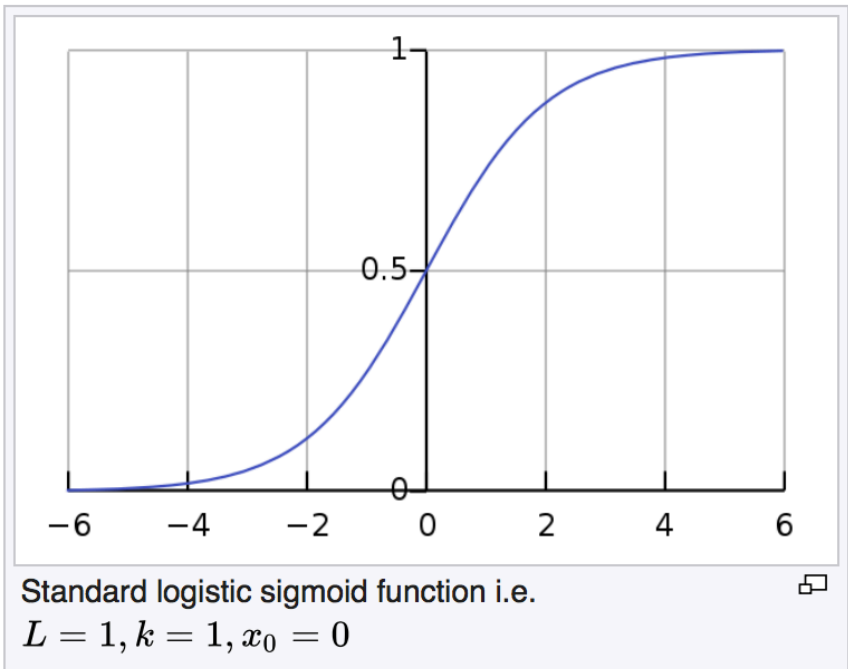
Logistic function

A **logistic function** or **logistic curve** is a common "S" shape (**sigmoid curve**), with equation:

$$f(x) = \frac{L}{1 + e^{-k(x-x_0)}}$$

where

- e = the **natural logarithm** base (also known as **Euler's number**),
- x_0 = the x -value of the sigmoid's midpoint,
- L = the curve's maximum value, and
- k = the steepness of the curve.^[1]

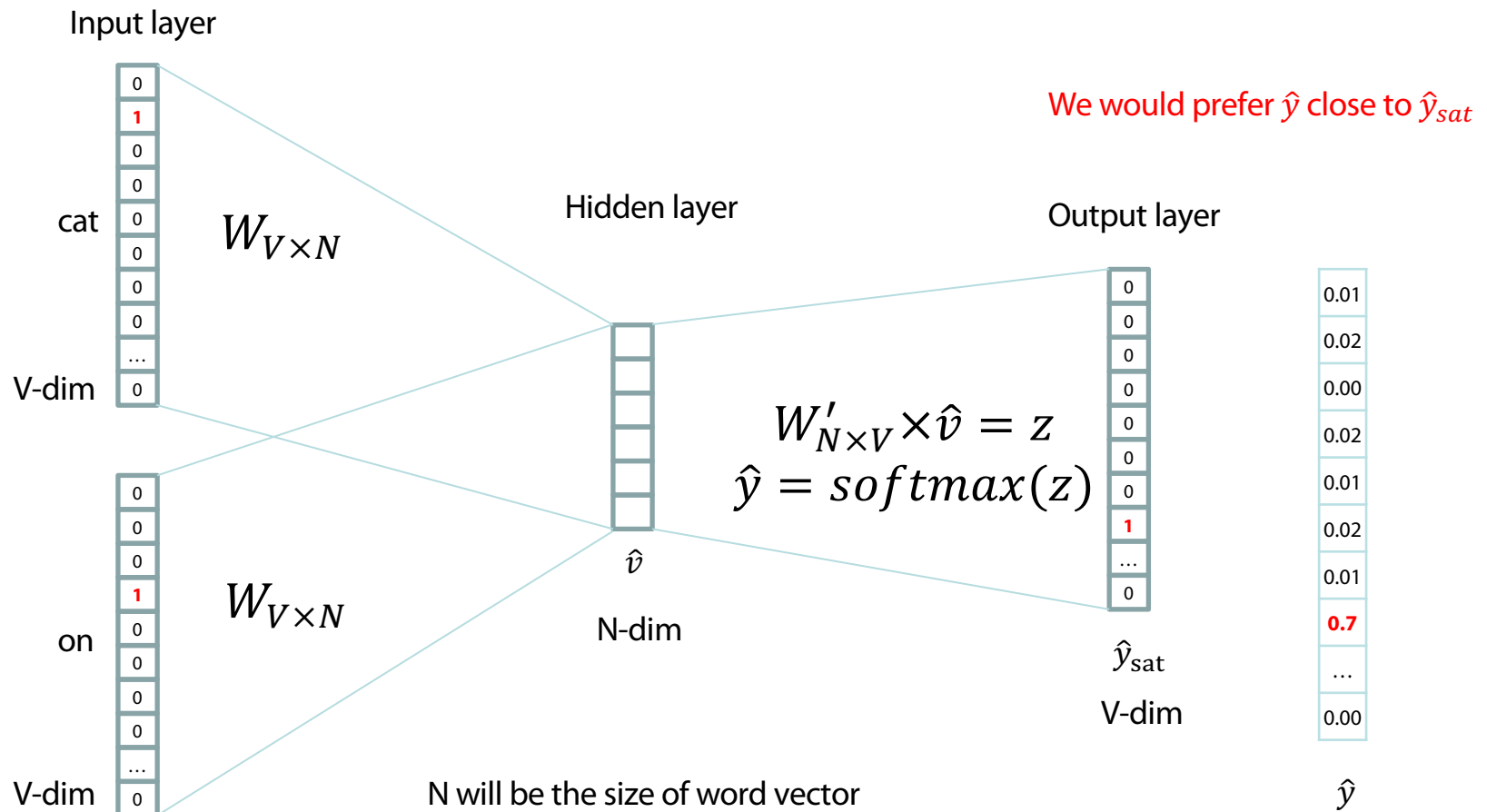


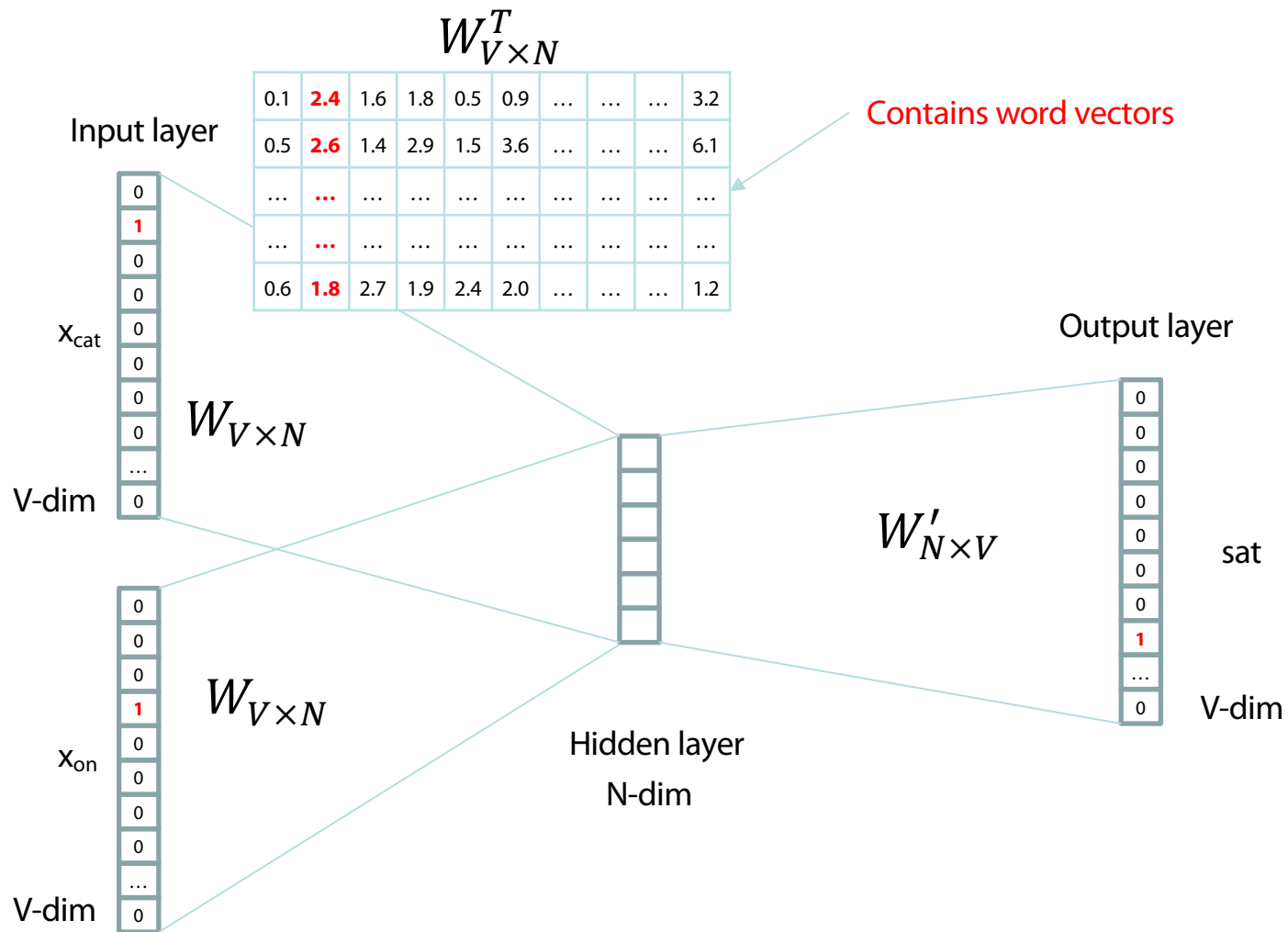
softmax(**z**)

The **softmax function**, or **normalized exponential function**, is a generalization of the **logistic function** that "squashes" a K -dimensional vector $\mathbf{z}$ of arbitrary real values to a K -dimensional vector $\sigma(\mathbf{z})$ of real values in the range $[0, 1]$ that add up to 1. The function is given by

$$\sigma : \mathbb{R}^K \rightarrow [0, 1]^K$$
$$\sigma(\mathbf{z})_j = \frac{e^{z_j}}{\sum_{k=1}^K e^{z_k}} \quad \text{for } j = 1, \dots, K.$$

In **probability theory**, the output of the softmax function can be used to represent a **categorical distribution** – that is, a **probability distribution** over K different possible outcomes.





Consider either W or W' as the word's representation.

Word Analogies

Test for linear relationships, examined by Mikolov et al. (2014)

a:b :: c:?



$$d = \arg \max_x \frac{(w_b - w_a + w_c)^T w_x}{\|w_b - w_a + w_c\| \|w_x\|}$$

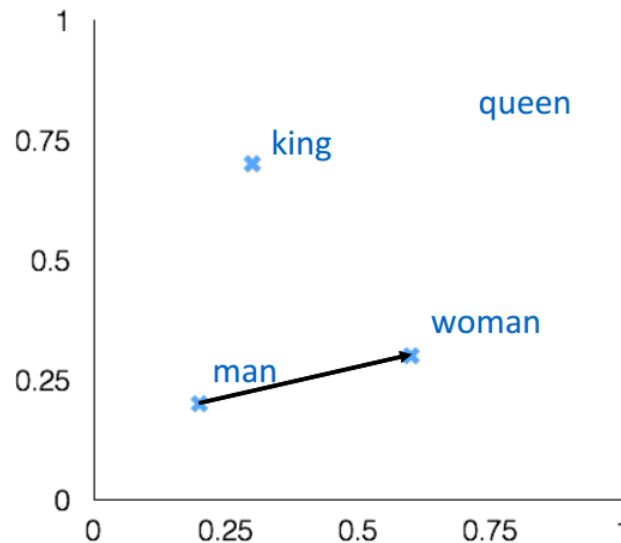
man:woman :: king:?

+ king [0.30 0.70]

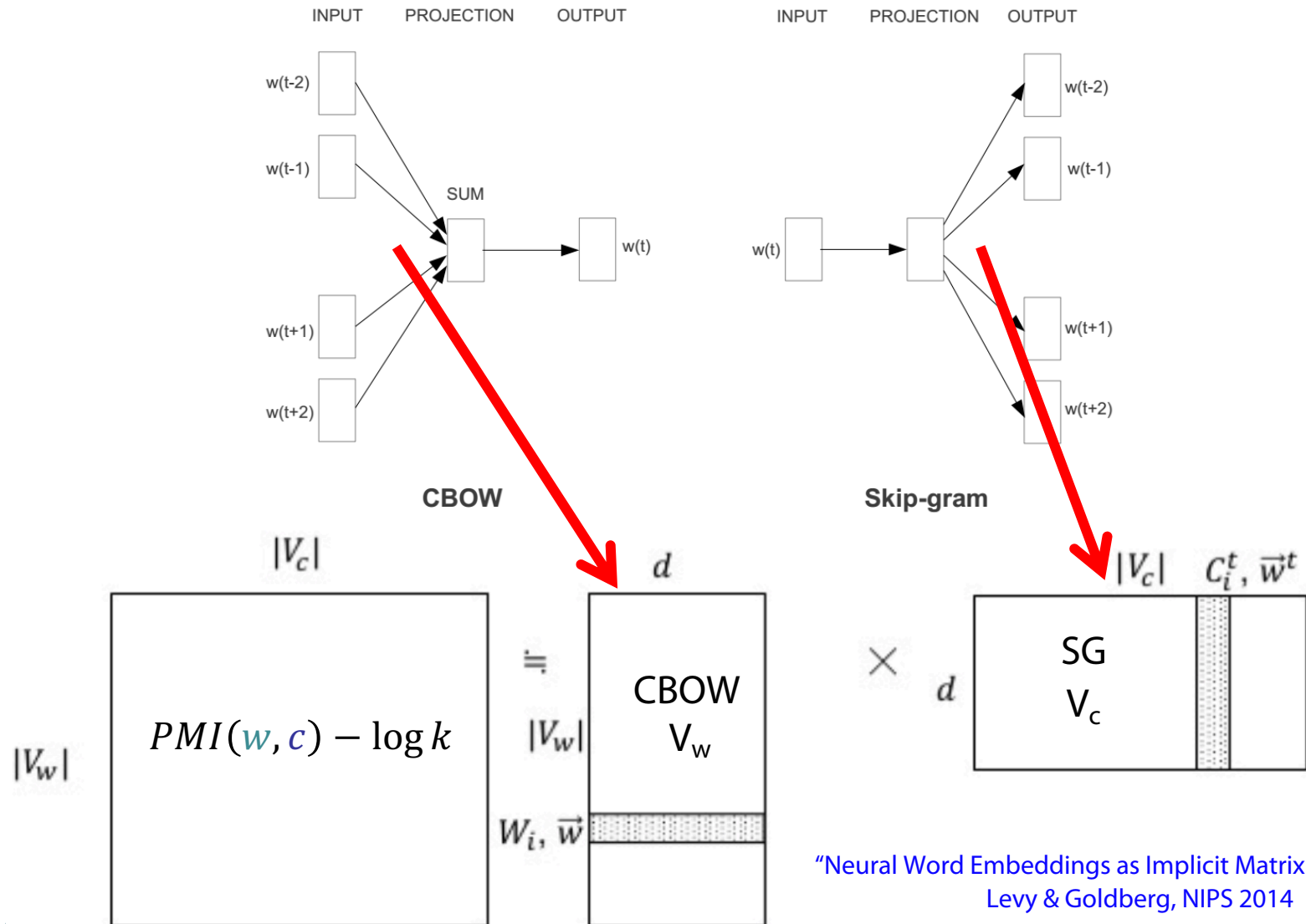
- man [0.20 0.20]

+ woman [0.60 0.30]

queen [0.70 0.80]



The Picture: CBOW and Skip-Gram (SG)



"Neural Word Embeddings as Implicit Matrix Factorization"
Levy & Goldberg, NIPS 2014

Convolution

Input image



Convolution
Kernel

$$\begin{bmatrix} -1 & -1 & -1 \\ -1 & 8 & -1 \\ -1 & -1 & -1 \end{bmatrix}$$

Feature map



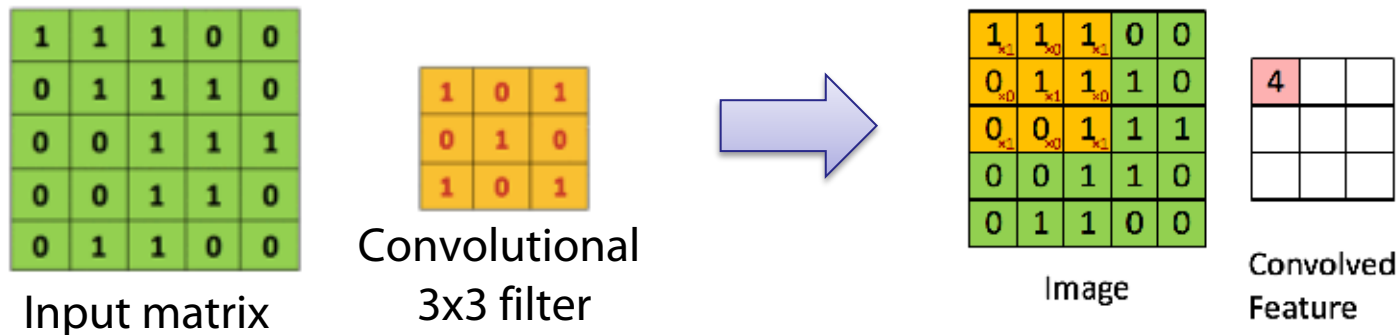
Convolutional Neural Networks (CNNs)

Main CNN idea for text:

Compute vectors for n-grams and group them afterwards

Example: “this takes too long” compute vectors for:

This takes, takes too, too long, this takes too, takes too long, this takes too long



http://deeplearning.stanford.edu/wiki/index.php/Feature_extraction_using_convolution

Convolutional Neural Networks (CNNs)

Main CNN idea for text:

Compute vectors for n-grams and **group them afterwards**

Feature Map

6	4	8	5
5	4	5	8
3	6	7	7
7	9	7	2

max pool
2x2 filters
and stride 2

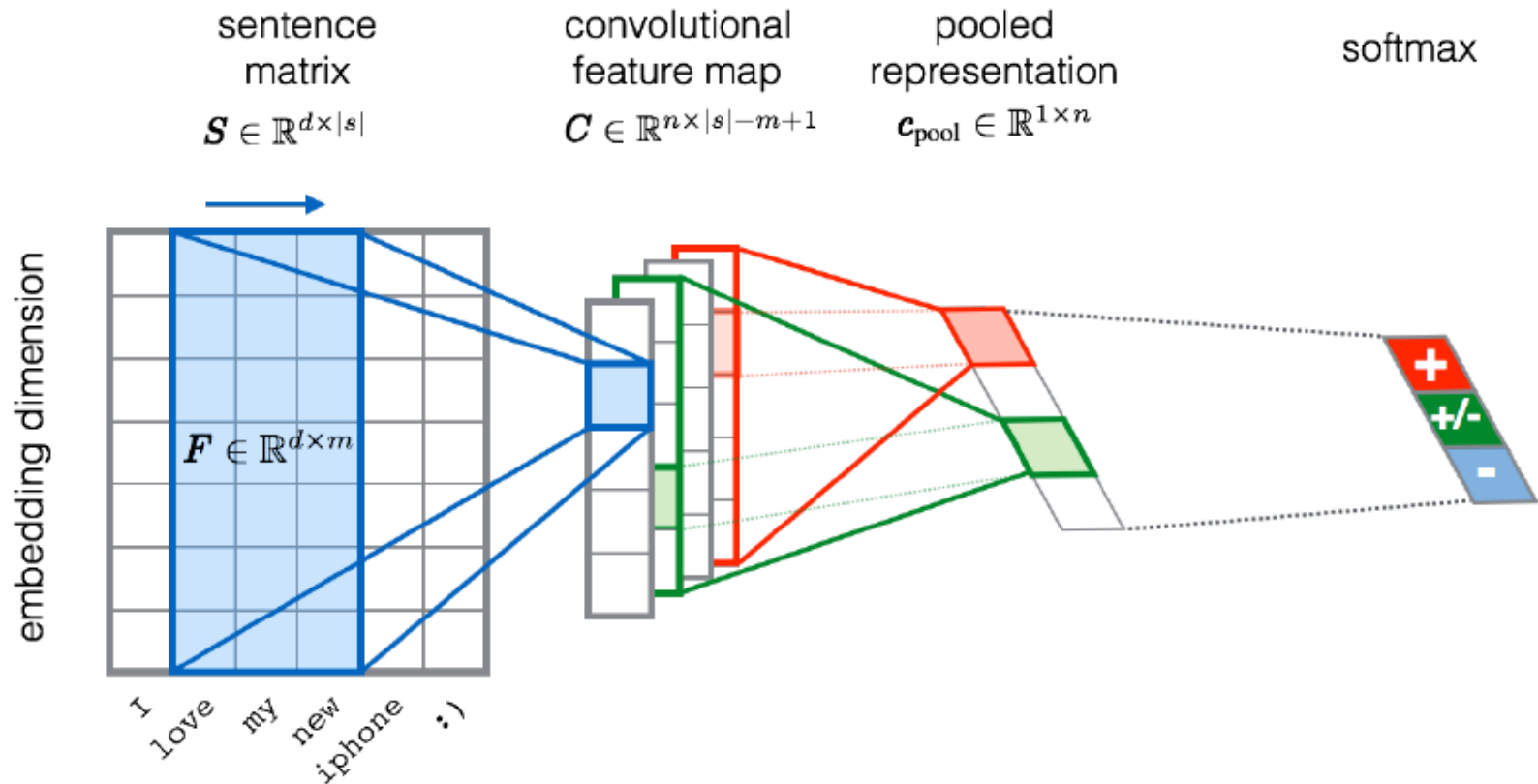


Max-Pooling

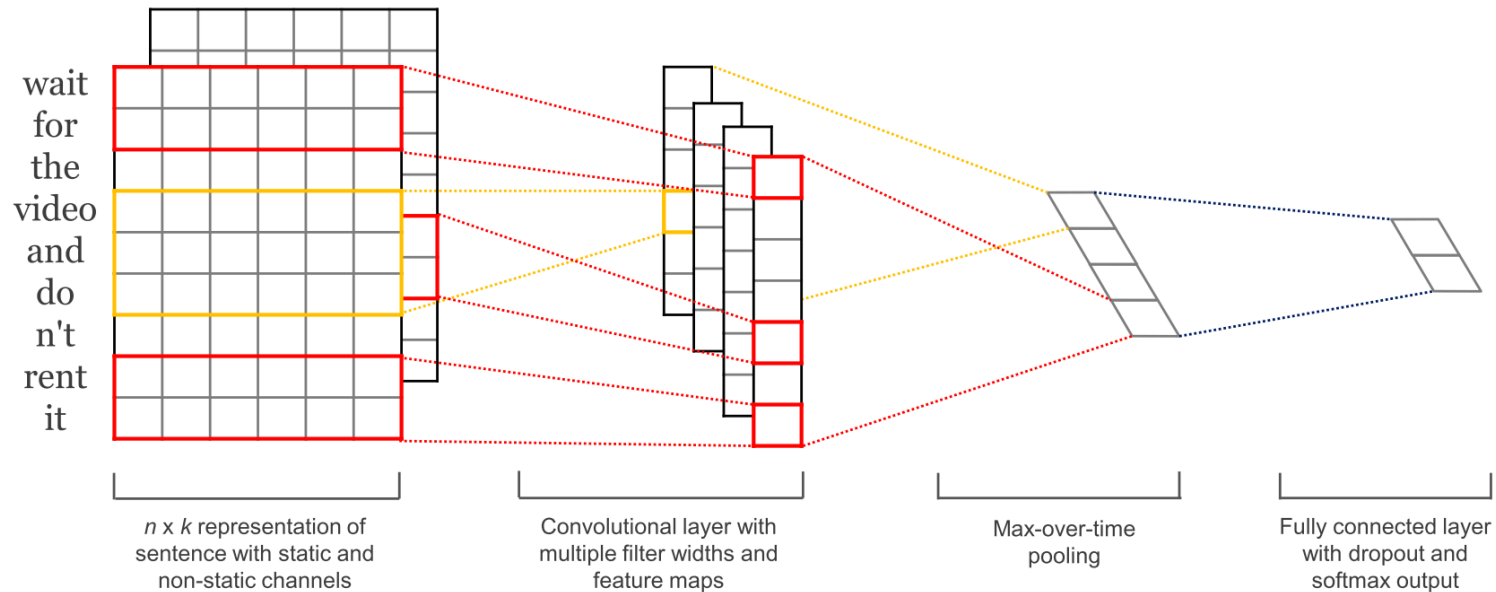
Dimensionsreduktion

<https://shafeentejani.github.io/assets/images/pooling.gif>

CNN for text classification



CNN with multiple filters



Kim, Y. "Convolutional Neural Networks for Sentence Classification", EMNLP (2014)

sliding over 3, 4 or 5 words at a time

Geschichtlicher Überblick

Anfänge

- **1943:** McCulloch und Pitts beschreiben und definieren eine Art erster neuronaler Netzwerke.
- **1949:** Formulierung der Hebb'schen Lernregel (nach Hebb)
- **1957:** Entwicklung des Perzeptrons durch Rosenblatt

Ernüchterung

- **1969:** Minsky und Papert untersuchen das Perzeptron mathematisch und zeigen dessen Grenzen, etwa beim XOR-Problem, auf.

Renaissance

- **1982:** Beschreibung der ersten selbstorganisierenden Netze (nach *biologischem Vorbild*) durch van der Malsburg und Kohonen und eines richtungweisenden Artikels von Hopfield, indem die ersten rückgekoppelten Netze (Hopfield-Netze, nach *physikalischen Vorbild*) beschrieben werden
- **1986:** Das Lernverfahren Backpropagation für mehrschichtige Perzeptrons wird entwickelt.

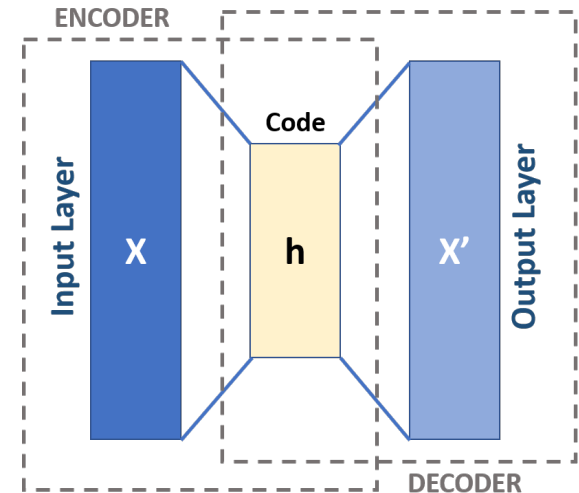
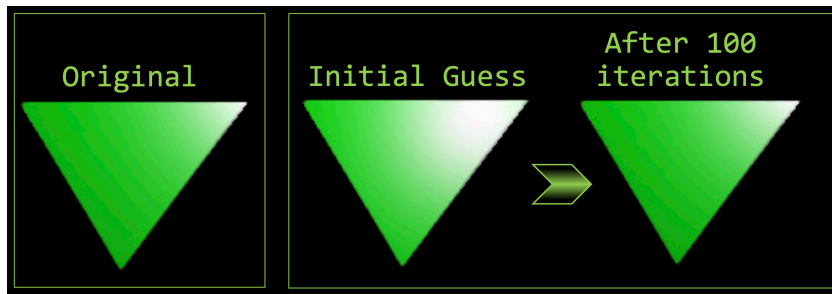
Boom

- **Ab 2000:** Deep Learning (Hinton, LeCun, Bengio, Ng, et al.)
- **Ab 2020:** Differentiable Programming (Lecun et al.)



Differential Programming: Basic Idea

- Example: Computer Vision
 - Estimate position of light source
 - Use standard learning approach



- Develop render in appropriate programming language (e.g., Julia)
- Differentiate renderer (e.g., Zygote)
 - Use differentiated render for backpropagation